



On active learning for Gaussian process-based global sensitivity analysis

Mohit S. Chauhan^a, Mariel Ojeda-Tuz^b, Ryan A. Catarelli^b, Kurtis R. Gurley^b,
Dimitrios Tsapetis^a, Michael D. Shields^{a,*}

^a Department of Civil and Systems Engineering, Johns Hopkins University, United States of America

^b Department of Civil and Coastal Engineering, University of Florida, United States of America

ARTICLE INFO

Keywords:

Sobol index
Active learning
Global sensitivity analysis
Gaussian process regression
Kriging

ABSTRACT

This paper explores the application of active learning strategies to adaptively learn Sobol indices for global sensitivity analysis. We demonstrate that active learning for Sobol indices poses unique challenges due to the definition of the Sobol index as a ratio of variances estimated from Gaussian process surrogates. Consequently, learning strategies must either focus on convergence in the numerator or the denominator of this ratio. However, rapid convergence in either one does not guarantee convergence in the Sobol index. We propose a novel strategy for active learning that focuses on resolving the main effects of the Gaussian process (associated with the numerator of the Sobol index) and compare this with existing strategies based on convergence in the total variance (the denominator of the Sobol index). The new strategy, implemented through a new learning function termed the MUSIC (minimize uncertainty in Sobol index convergence), generally converges in Sobol index error more rapidly than the existing strategies based on the Expected Improvement for Global Fit (EIGF) and the Variance Improvement for Global Fit (VIGF). Both strategies are compared with simple sequential random sampling and the MUSIC learning function generally converges most rapidly for low-dimensional problems. However, for high-dimensional problems, the performance is comparable to random sampling. The new learning strategy is demonstrated for a practical case of adaptive experimental design for large-scale Boundary Layer Wind Tunnel experiments.

1. Introduction

Active learning is widely used in computational physics-based modeling, and increasingly in experimental studies [1], to serve a variety of objectives. In these schemes, a machine learning (ML) model is employed in an iterative routine to extract information from previous simulations/experiments and identify, through some *learning function* the most informative simulation/experiment to run next. This has been especially beneficial in the fields of optimization and reliability, although it has also found wider applications in uncertainty quantification (UQ) more generally.

In optimization, active learning is at the heart of Bayesian Optimization (BO). In BO, Bayes Rule is used to define a learning function (termed an acquisition function in BO) that seeks to identify the minimum (or maximum) of some function. Comprehensive reviews of BO can be found in the following Refs. [2–4]. A detailed discussion of BO is beyond the scope of this work, but the most common approach in these methods is to approximate the function with a Gaussian Process (GP) and construct a learning function that exploits the jointly Gaussian nature of the function to adaptively learn where the sought extremum has the highest probability of lying. Among these developments, the

most widely used learning functions are the Expected Improvement Function (EIF) [5–7] the Knowledge Gradient [8,9] and the Entropy Search [10,11] algorithms.

For reliability analysis, the objective is to estimate the probability of failure of a physical system, which involves solving a high-dimensional probability integral (expectation) over a binary classifier that indicates either safety or failure. This integral is often solved statistically using Monte Carlo methods wherein the integral is estimated as the statistical expectation of the failure indicator by sampling from the input distribution and assessing the indicator at each sample point. Active learning is then used to approximate this indicator function efficiently to minimize the computational burden associated with assessing system failure at each sample with an expensive computational model. Numerous learning functions have been developed for this task, with each attempting to efficiently approximate the limit surface (surface that separates safe and failure domains) in different ways, but in most cases using a GP. The first such method, termed the Expected Feasibility Function (EFF) was proposed by Bichon et al. in 2008 [12]. This was followed by the widely used U-function developed by Echard et al. [13] in the

* Corresponding author.

E-mail address: michael.shields@jhu.edu (M.D. Shields).

Adaptive Kriging with Monte Carlo Simulation (AK-MCS) method. Since then, various modifications have been made and these concepts have been integrated with other reliability methods (such as subset simulation [14,15]) to improve the efficiency, convergence, and accuracy of these methods [16–20].

Although optimization and reliability are the most common application domains for active learning, it has been applied more broadly in UQ, for example to efficiently construct surrogate models (aka metamodels or emulators) for expensive physics-based computational models. Lam [21], for example, proposed the Expected Improvement for Global Fit (EIGF) learning function to improve the convergence of GP surrogate models for UQ purposes. Active learning has also been applied using other ML surrogate model forms, such as Polynomial Chaos Expansions (PCE) [22–24] and Deep Neural Networks (DNNs) [25–27]. Furthermore, recent efforts have integrated active learning concepts into the design of physical experiments for wind tunnel experiments [1,28–30], materials discovery [31,32], and biological sciences [33,34].

Regardless of their application or objective, a common challenge in active learning is to balance *exploration* and *exploitation*. That is, any active learning method (i.e. learning function) must simultaneously use (exploit) the information obtained from prior simulations/experiments to guide further analyses and, at the same time, explore new regions of the parameter space where little or no data has been collected and uncertainty is correspondingly high. Balancing these objectives is a primary challenge in active learning [35], especially in high-dimensional parameter spaces where exploration can be particularly difficult [36].

In this work, we explore the potential for active learning for Global Sensitivity Analysis (GSA), which to the authors' knowledge, has not previously been studied. Active learning for GSA is of particular interest here because of its potential to aid in other UQ studies. In particular, GSA can be used as an initial filter in UQ studies to reduce the number of random variables that need to be considered in the study. For example, in Section 6 we consider a practical boundary layer wind tunnel (BLWT) experimental design problem that has 10 input random variables and, through active learning, determine that only 3 random variables contribute appreciably to the variance of the response. We can then practically ignore uncertainty in the other 7 random variables, thus making UQ more computationally tractable in this complex physical setting. Hence, active learning for GSA allows us to adaptively reduce the dimension rapidly *on-the-fly* and therefore speed up the overall UQ task.

We show that active learning for GSA presents some unique challenges that do not arise in other active learning domains. These challenges arise from the definition of Global Sensitivity Indices (specifically Sobol indices [37]) as the ratio of two variances, which makes the construction of a learning function that directly targets the Sobol index difficult to derive. Instead, we explore two different learning functions — one that targets efficient learning for the numerator and one that targets the denominator. We show that each can be beneficial under certain circumstances, but neither approach universally outperforms random sampling. We explore the reasons for this, and specifically demonstrate that this results from the challenges of minimizing the error in a ratio. We discuss when the different approaches appear to be beneficial using a series of analytical functions of varying dimensions. We then demonstrate our use of active learning for GSA on the BLWT application of interest.

2. Preliminaries

Here, we briefly review the foundations of Global Sensitivity Analysis using Sobol Indices and Gaussian Process (GP) regression, upon which all further analyses will depend.

2.1. Sobol sensitivity indices

Variance-based methods for GSA aim to decompose the variance of the model output $y(\mathbf{X})$ into distinct contributions from each of the individual random variables X_i (main effects) and the interaction of multiple random variables. In the Sobol method [37–39], this is done by first constructing the high-dimensional model representation (HDMR) of the model as:

$$y(\mathbf{X}) = f_0 + \sum_{i=1}^d f_i(X_i) + \sum_{i=1}^d \sum_{j>i}^d f_{ij}(X_i, X_j) + \dots \quad (1)$$

where

$$\begin{aligned} f_0 &= E[Y] \\ f_i &= E[Y|X_i] - E[Y] \\ f_{ij} &= E[Y|X_i, X_j] - f_i - f_j - E[Y] \end{aligned} \quad (2)$$

Taking the variance of Eq. (1) and exploiting the orthogonality between terms in the expansion yields the following variance decomposition:

$$\text{Var}(Y) = \sum_{i=1}^d V_i + \sum_{i=1}^d \sum_{j>i}^d V_{ij} + \dots \quad (3)$$

where V_i represent the main effect variances and V_{ij} represent the variance contribution from interactions between variables X_i and X_j given by:

$$\begin{aligned} V_i &= \text{Var}_{X_i}(E_{X_{\sim i}}[Y|X_i]) \\ V_{ij} &= \text{Var}_{X_{ij}}(E_{X_{\sim ij}}[Y|X_i, X_j]) - \text{Var}_{X_i}(E_{X_{\sim i}}[Y|X_i]) - \text{Var}_{X_j}(E_{X_{\sim j}}[Y|X_j]) \end{aligned} \quad (4)$$

The Sobol sensitivity indices reflect the relative contribution of each term in Eq. (4) to the total variance and are given by:

$$S_i = \frac{V_i}{\text{Var}(Y)} = \frac{\text{Var}_{X_i}(E_{X_{\sim i}}[Y|X_i])}{\text{Var}(Y)} \quad (5)$$

$$\begin{aligned} S_{ij} &= \frac{V_{ij}}{\text{Var}(Y)} \\ &= \frac{\text{Var}_{X_{ij}}(E_{X_{\sim ij}}[Y|X_i, X_j]) - \text{Var}_{X_i}(E_{X_{\sim i}}[Y|X_i]) - \text{Var}_{X_j}(E_{X_{\sim j}}[Y|X_j])}{\text{Var}(Y)} \end{aligned} \quad (6)$$

where S_i are the main effect sensitivity indices and S_{ij} are the interaction sensitivities and we have that:

$$\sum_{i=1}^d S_i + \sum_{i=1}^d \sum_{j>i}^d S_{ij} + \dots = 1. \quad (7)$$

Here, we will focus predominantly on the main effect sensitivity indices in Eq. (5), with some supporting discussion of interaction sensitivities. Importantly, we highlight that the sensitivity index S_i is defined as a ratio of variances (the so-called main effect variance in the numerator and total variance in the denominator), which *both* must be estimated in an active learning framework.

2.2. Gaussian process regression/Kriging

Gaussian Process (GP) regression is a machine learning (ML) method that seeks to identify the best fit Gaussian stochastic process to a set of data points [40]. Kriging [41–43], a variation of GP regression, is a special case wherein the GP regressor serves to interpolate between the set of training points. Kriging is a widely-used ML method to construct approximate surrogate models for complex computational models. It is particularly attractive for active learning because it provides a model predictor as well as a measure of the uncertainty in the prediction of the Gaussian posterior standard deviation.

Consider a computational model $y(\mathbf{x})$ having input vector $\mathbf{x} = \{x_1, x_2, \dots, x_d\} \in \mathbb{R}^d$. Next, consider that we have n realizations of $\mathbf{x}$,

denoted by $\mathbf{X} = \{\mathbf{x}^{(1)}, \mathbf{x}^{(2)}, \dots, \mathbf{x}^{(n)}\}$ and corresponding model outputs $\mathbf{Y} = \{y^{(1)}, y^{(2)}, \dots, y^{(n)}\}$ evaluated by $y^{(i)} = y(\mathbf{x}^{(i)})$. The pairs, $\mathbf{X}, \mathbf{Y}$ serve as the training data for the Gaussian process regression wherein the model $y(\mathbf{x})$ is approximated by:

$$\mathcal{Y}(\mathbf{x}, \omega) = F(\mathbf{x}) + Z(\mathbf{x}, \omega) \quad (8)$$

where $F(\cdot)$ is a regression model and $Z(\cdot)$ is a zero-mean Gaussian random process having sample space indexed by $\omega \in \Omega$. We consider that the regression model $F(\cdot)$ is defined through a linear combination of basis functions $\mathbf{f}(\mathbf{x}) = \{f_1(\mathbf{x}), f_2(\mathbf{x}), \dots, f_p(\mathbf{x})\}$ having coefficients $\boldsymbol{\beta} = \{\beta_1, \beta_2, \dots, \beta_p\}$ as:

$$F(\mathbf{x}) = \boldsymbol{\beta}^T \mathbf{f}(\mathbf{x}). \quad (9)$$

The Gaussian random process $Z(\cdot)$ is considered to have zero mean and covariance:

$$\mathbb{E}[Z(\mathbf{x}_1)Z(\mathbf{x}_2)] = \sigma_z^2 \mathcal{R}(\mathbf{x}_1, \mathbf{x}_2 | \boldsymbol{\theta}) \quad (10)$$

where $\mathcal{R}(\cdot)$ is the autocorrelation function having hyperparameters $\boldsymbol{\theta}$. In this work, we will consider the Gaussian correlation model having form:

$$\mathcal{R}(\mathbf{x}_1, \mathbf{x}_2 | \boldsymbol{\theta}) = \prod_{k=1}^d \exp(-\theta_k |x_{2k} - x_{1k}|^2) = \prod_{k=1}^d r_k(x_{1k}, x_{2k} | \theta_k). \quad (11)$$

The GP is then fit to the normalized training data by identifying the posterior distribution of the random process given the data. Given the Gaussian assumption, the posterior distribution is determined by considering that the joint distribution of the GP predictions $\mathcal{Y}(\mathbf{x})$ and the observations $\mathbf{Y}$ is a multivariate Gaussian and can be expressed as:

$$\begin{Bmatrix} \mathcal{Y}(\mathbf{x}) \\ \mathbf{Y} \end{Bmatrix} \sim N \left(\begin{Bmatrix} \mathbf{f}(\mathbf{x})^T \boldsymbol{\beta} \\ \mathbf{F} \boldsymbol{\beta} \end{Bmatrix}, \sigma_z^2 \begin{Bmatrix} 1 & \mathbf{r}^T(\mathbf{x}) \\ \mathbf{r}(\mathbf{x}) & \mathbf{R} \end{Bmatrix} \right) \quad (12)$$

where $\mathbf{F}$ is the matrix of basis function evaluations at the training points given by $F_{ij} = f_j(\mathbf{x}^{(i)})$, $i = 1, \dots, n$, $j = 1 \dots p$, $\mathbf{r}(\mathbf{x})$ is the vector of correlations between the prediction point $\mathbf{x}$ and the training points $\mathbf{x}^{(i)}$ given by $r_i = \mathcal{R}(\mathbf{x}, \mathbf{x}^{(i)} | \boldsymbol{\theta})$, $i = 1, \dots, n$, and $\mathbf{R}$ is the correlation matrix of points in the training set given by $R_{ij} = \mathcal{R}(\mathbf{x}^{(i)}, \mathbf{x}^{(j)} | \boldsymbol{\theta})$, $i, j = 1, \dots, n$.

This posterior can be estimated by identifying the hyperparameters $\boldsymbol{\theta}$, variance σ_z^2 , and the regression coefficients $\boldsymbol{\beta}$. This is typically done by solving for the parameters that either maximize the likelihood of the observations or minimize the cross-validation error. We apply maximum likelihood estimation where the likelihood function follows from the Gaussian assumption as:

$$\mathcal{L}(\boldsymbol{\theta}, \sigma_z^2, \boldsymbol{\beta} | \mathbf{Y}) = \frac{(\det \mathbf{R})^{-1/2}}{(2\pi\sigma_z^2)^{n/2}} \exp \left[-\frac{1}{2\sigma_z^2} (\mathbf{Y} - \mathbf{F} \boldsymbol{\beta})^T \mathbf{R}^{-1} (\mathbf{Y} - \mathbf{F} \boldsymbol{\beta}) \right] \quad (13)$$

After estimating the hyperparameters, the mean and the variance of the Gaussian random variable conditioned on the training data are given by:

$$\hat{y}(\mathbf{x}) = \mathbf{f}(\mathbf{x})^T \boldsymbol{\beta} + \mathbf{r}(\mathbf{x})^T \mathbf{R}^{-1} (\mathbf{Y} - \mathbf{F} \boldsymbol{\beta}) \quad (14)$$

$$\sigma_y^2(\mathbf{x}) = \sigma_z^2 (1 - \mathbf{r}(\mathbf{x})^T \mathbf{R}^{-1} \mathbf{r}(\mathbf{x}) + \mathbf{t}(\mathbf{x})^T (\mathbf{F}^T \mathbf{R}^{-1} \mathbf{F})^{-1} \mathbf{t}(\mathbf{x})) \quad (15)$$

where

$$\mathbf{t}(\mathbf{x}) = \mathbf{F}^T \mathbf{R}^{-1} \mathbf{r}(\mathbf{x}) - \mathbf{f}(\mathbf{x}) \quad (16)$$

These provide the Kriging prediction at a new point $\hat{y}(\mathbf{x})$ along with a measure of its uncertainty $\sigma_y(\mathbf{x})$.

3. Estimating Sobol indices from Gaussian process regression

In this section, we review the process through which Sobol sensitivity indices can be estimated from a GP regression model as originally derived in [44]. Marrel et al. [44] presented two approaches to compute Sobol indices using GPs. In the first approach, the Sobol indices are

computed simply using the GP predictor, in which case the method provides a deterministic point estimator based purely on the conditional mean as:

$$S_i = \frac{\text{Var}_{X_i}(E_{X_{\sim i}}[E_{\Omega}[\mathcal{Y}(X, \omega) | X_i]])}{\text{Var}_X(E_{\Omega}[\mathcal{Y}(X, \omega)])} = \frac{\text{Var}_{X_i}(E_{X_{\sim i}}[\hat{y}(X) | X_i])}{\text{Var}_X(\hat{y}(X))} \quad (17)$$

where $\hat{y}(X)$ denotes the conditional mean in Eq. (14) and is treated as a deterministic function in this approach.

The second approach, which we will employ in this study, uses the full GP model as a stochastic function, which results in random variable sensitivity indices given by:

$$\tilde{S}_i(\omega) = \frac{\text{Var}_{X_i}(E_{X_{\sim i}}[\mathcal{Y}(X, \omega) | X_i])}{E_{\Omega}[\text{Var}_X(\mathcal{Y}(X, \omega))]} \quad (18)$$

The sensitivity measure ($\tilde{S}_i(\omega)$) is obtained using the global GP, $\mathcal{Y}(X, \omega)$, where the denominator considers the complete variance of the GP itself over all X , which remains a random variable indexed on ω . The expectation in the denominator therefore reduces this to a deterministic expected variance. Further, the mean and variance of the Sobol indices are given by:

$$\mu_{\tilde{S}_i} = \frac{E_{\Omega}[\text{Var}_{X_i}(E_{X_{\sim i}}[\mathcal{Y}(X, \omega) | X_i])]}{E_{\Omega}[\text{Var}_X(\mathcal{Y}(X, \omega))]} \quad (19)$$

and

$$\sigma_{\tilde{S}_i}^2 = \frac{\text{Var}_{\Omega}(\text{Var}_{X_i}(E_{X_{\sim i}}[\mathcal{Y}(X, \omega) | X_i])]}{(E_{\Omega}[\text{Var}_X(\mathcal{Y}(X, \omega))])^2} \quad (20)$$

where $\mu_{\tilde{S}_i}$ provides an estimate of the Sobol index and $\sigma_{\tilde{S}_i}^2$ gives an estimate of its uncertainty. To compute the first-order Sobol indices, we start by taking the expectation of the GP $\mathcal{Y}(X, \omega)$ over all the inputs except ' X_i ' in the numerator of Eq. (18) and denote this as:

$$A_i(X_i, \omega) = E_{X_{\sim i}}[\mathcal{Y}(X, \omega) | X_i]. \quad (21)$$

Since, $\mathcal{Y}(X, \omega)$ is a Gaussian process and expectation is a linear operator, $A_i(X_i, \omega)$ is also a GP referred to as the main effect GP for i th input dimension. Hereafter $A_i(X_i, \omega)$ is denoted as $A_i(X_i)$ for brevity in notation.

The mean and covariance function of the main effect GP can be determined by integrating the original GP with respect to the joint probability measure over all inputs except X_i . Considering independent inputs, the mean function is given by [45]:

$$\mu_{A_i}(x_i) = E_{X_{\sim i}}[A_i(X_i)] = \int_{\mathbf{x}_{\sim i}} \hat{y}(\mathbf{x}) \prod_{j \neq i} p_{X_j}(x_j) d\mathbf{x}_j \quad (22)$$

which is easily computed as a product of one-dimensional integrals. Again considering independent inputs, the covariance function of the main effect GP is given by [45]:

$$\begin{aligned} \text{Cov}_{X_{\sim i}, X_{\sim i}}(A_i(X_{1i}), A_i(X_{2i})) \\ = \int_{\mathbf{x}_{1 \sim i}} \int_{\mathbf{x}_{2 \sim i}} \text{Cov}(\mathcal{Y}(\mathbf{x}_1), \mathcal{Y}(\mathbf{x}_2)) \prod_{j \neq i} p_{X_j}(x_{1j}) d\mathbf{x}_{1j} \prod_{k \neq i} p_{X_k}(x_{2k}) d\mathbf{x}_{2k} \end{aligned} \quad (23)$$

which can be expressed as a product of 2-dimensional integrals. Detailed derivations for computing the mean and covariance functions of the main effect GPs in closed form under specific conditions are provided in Appendix A. We further note that Eqs. (22) and (23) can be generalized for interactions between variables as shown in Appendix B. Computation of the interaction sensitivities follows directly from the following for main effect sensitivities. These interaction sensitivities can, in principle, be used within the learning although this is outside the scope of the present work.

To compute Sobol indices, we need to compute the variance of the main effect GP as:

$$\sigma_{A_i}^2 = \int_{X_i} (A_i(x_i) - E[A_i(X_i)])^2 p_{X_i}(x_i) dx_i \quad (24)$$

Here, a Monte Carlo estimate of this integral is computed by defining

a random discretization of the sample space of X_i as $\{a_1, a_2, \dots, a_{n_j}\}$ (note these points will be reused later for learning function evaluations as well) and constructing the following Gaussian random vector to approximate the main effect GP:

$$V_i = [A_i(a_1), A_i(a_2), \dots, A_i(a_{n_j})]^T$$

We can then develop the following two Monte Carlo estimators of the Sobol indices.

3.1. SI computation using the mean of the main effect GP

In the first approach, the mean of vector V_i (i.e. $\mu_{V_i} = E_{X_i}[V_i]$) is computed using Eqs. (22), where each element of the mean vector is defined as $\mu_{V_{i,j}} = \mu_{A_i}(a_j) = E[A_i(a_j)] \forall j \in \{1, 2, \dots, n_j\}$. Then, the following standard Monte Carlo estimators are used to estimate the variance of the main effect GP:

$$\hat{\sigma}_{A_i}^2 = \frac{1}{n_j - 1} \sum_{j=1}^{n_j} (\mu_{V_{i,j}} - \bar{\mu}_{V_i})^2 \quad (25)$$

where $\bar{\mu}_{V_i} = \frac{1}{n_j} \sum_{j=1}^{n_j} \mu_{V_{i,j}}$ is the average of the mean vector. The first-order Sobol indices are then estimated by:

$$\mu_{S_i} = \frac{\hat{\sigma}_{A_i}^2}{\hat{\sigma}^2} \quad \forall i \in \{1, 2, \dots, d\} \quad (26)$$

where $\hat{\sigma}^2$ is the total variance of $\mathcal{Y}(x)$. This estimator provides the Sobol index as identified from the mean GP predictor for the main effect, but does not account for uncertainty/variability in the main effect GP. In the next section, an estimator for the standard deviation of the Sobol index is also given.

3.2. SI computation using the main effect GP

Here, we simulate the main effect GP and use these simulations to estimate both the mean and the standard deviation of the Sobol index. First, the mean vector and covariance matrix of V_i are computed using Eqs. (22) and (23), respectively. We then generate n_s realizations of the random vector V_i using the Cholesky decomposition of the covariance matrix [42], $C = LL^T$, where the k th realization for i th input (i.e. V_{ik}) is given as follows:

$$V_{ik} = E[A_i(X_i)] + L\epsilon_k = [A_{ik}(a_1), A_{ik}(a_2), \dots, A_{ik}(a_{n_j})]^T \quad (27)$$

where ϵ_k is a realization of an n_j dimensional zero-mean, uncorrelated Gaussian random vector. Then, the following standard Monte Carlo estimators are used to estimate the variance of the main effect GP:

$$\hat{\sigma}_{A_i}^2 = \frac{1}{n_s} \sum_{k=1}^{n_s} \hat{\sigma}_{A_{ik}}^2 = \frac{1}{n_s} \sum_{k=1}^{n_s} \frac{1}{n_j - 1} \sum_{j=1}^{n_j} (A_{ik}(a_j) - \hat{\mu}_{V_{ik}})^2 \quad (28)$$

where $\hat{\mu}_{V_{ik}}$ is the sample mean for vector V_{ik} , $\hat{\sigma}_{A_{ik}}^2$ is the k th estimate of the variance of main effect GP (A_i), and the corresponding standard error is given by:

$$\hat{S}_{A_i}^2 = \frac{1}{n_s - 1} \sum_{k=1}^{n_s} (\hat{\sigma}_{A_{ik}}^2 - \hat{\sigma}_{A_i}^2)^2. \quad (29)$$

From these, the mean and standard deviation of the first-order Sobol indices are estimated as:

$$\mu_{S_i} = \frac{\hat{\sigma}_{A_i}^2}{\hat{\sigma}^2} \quad \text{and} \quad \sigma_{S_i}^2 = \frac{\hat{S}_{A_i}^2}{\hat{\sigma}^2} \quad \forall i \in \{1, 2, \dots, d\} \quad (30)$$

where $\hat{\sigma}^2$ is the total variance of $\mathcal{Y}(x)$.

3.3. Computational considerations

The computational cost of applying these two methods for Sobol index estimation are vastly different and can play a critical role when integrated into an active learning loop where estimates need to be computed repeatedly. Moreover, we have to repeat these estimates for each input dimension, which can be problematic for problems with a large number of inputs. This must be weighed against the accuracy of the two approaches and the need to estimate the standard deviation of the Sobol index. The first approach only utilizes the mean predictor of Main Effect GP, which is computationally inexpensive but does not estimate the uncertainty in Sobol indices. The second approach considers the complete probability structure of the Main Effect GP and therefore provides an error measure for each Sobol index. But it comes at a high computational cost because it requires a large set of sample points (n_j should be large) to compute Sobol indices with high accuracy. This results in large covariance matrices, which can make the main effect simulations in Eq. (27) expensive. We therefore employ the first approach to estimate Sobol indices because we do not require uncertainty estimates on the Sobol indices at each iteration.

4. Active learning for Sobol Index estimation

Next, we compare strategies to perform active learning for Sobol Index estimation. We specifically consider two different learning functions designed with the GP-based Sobol index estimates from Eq. (18) in mind. We first recognize that the Sobol indices given by Eq. (18) are expressed as a ratio of two variances. The numerator is the variance of the main effect GP and can be expressed as $\sigma_{A_i}^2 = \text{Var}_{X_i}(A_i(X_i))$, while the denominator is the expected value of the variance of the complete GP, $\mathcal{Y}(X, \omega)$. The challenge of active learning in this case is therefore to minimize uncertainty in both components. To do so is not straightforward. We therefore propose two learning strategies that aim to minimize uncertainty in either the numerator or the denominator.

According to the first strategy, we aim to minimize uncertainty in the overall GP surrogate (specifically output variance, the Sobol index denominator). Several learning functions have been proposed to achieve this objective including the Expected Improvement for Global Fit (EIGF — described in Section 4.1 [21]), Variance Improvement for Global Fit (VIGF — described in Section 4.2) [46], Integrated Variance Reduction (IVR) [47], and Posterior Variance Contribution (PVC) [48]. In this study, we specifically select the EIGF and VIGF as exemplars for this class of methods. Note that the PVC has specifically been used to compute Sobol' indices [49], although the learning function itself does not explicitly account for the main effects necessary for their computation. For the second strategy, we develop a new learning function, referred to as the MUSIC (Minimize Uncertainty in Sobol Index Computation) learning function, designed to select samples that reduce uncertainty in the main effect GPs $A_i(X_i)$ and described in Section 4.3.

These two learning strategies are integrated into a standard active learning loop illustrated in Fig. 1, which operates as follows. First, a set of initial training data (X, Y) are generated by randomly sampling the input random vector x and evaluating the computational model $y(x)$ for each sample. Then, an initial GP surrogate model is fit to the data and initial estimates of the Sobol indices can be made using the estimators above. Next, we generate a large set of n_j candidate samples (X^c of the input random variables (typically $n_j \geq 10000$) and evaluate the appropriate learning function at each candidate sample. According to the specified learning function, a candidate sample, say x^* , is selected and the model is evaluated at this point, $y(x^*)$. A new GP is fit to the augmented data set and the Sobol indices are updated. A convergence criterion on the Sobol indices is assessed. If the Sobol indices are deemed sufficiently accurate, the process stops. If the indices are not considered converged, then a new set of candidate points are generated and the iterations continue.

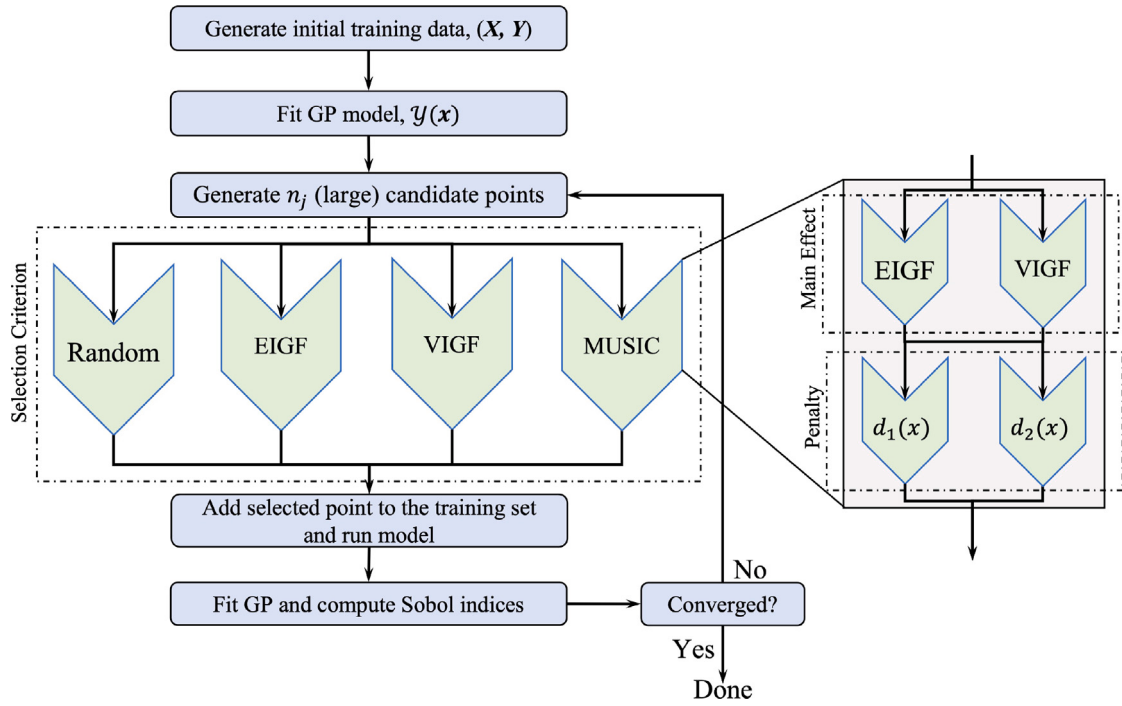


Fig. 1. Flowchart of the active learning strategies for global sensitivity analysis. Each of the learning functions (and random sampling) shown in green are studied in this work.

We compare both of these strategies with a simple random sampling approach in which, rather than employing a learning function in each iteration, one of the n_j candidate samples is selected at random. This provides a baseline to measure whether the active learning strategies are effective for improving the convergence of Sobol index estimates.

4.1. Expected improvement for global fit (EIGF)

To improve the fit of the global GP, $\mathcal{Y}(X, \omega)$, and specifically minimize uncertainty in the estimate of its variance, we employ the EIGF. Proposed by Lam [21], the EIGF defines an improvement function given by:

$$I(x) = (\mathcal{Y}(x) - y(x^{(j^*)}))^2 \quad (31)$$

where $y(x^{(j^*)})$ is the observation of y at the training point $x^{(j^*)}$ ($x^{(j^*)} \in X$) that is nearest to the candidate point x ($x \in X^c$) and $\mathcal{Y}(x)$ is the normal random variable resulting from evaluating the GP at point x . Taking the expectation of Eq. (31) yields:

$$\text{EIGF}(x) = E[I(x)] = (\hat{y}(x) - y(x^{(j^*)}))^2 + \sigma_y^2(x) \quad (32)$$

where, $\hat{y}(x)$ and $\sigma_y^2(x)$ are the GP mean and variance at point x (see Eqs. (14) and (15)). The EIGF learning function aims to identify the point x^* that maximizes the expectation in Eq. (32) as:

$$x^* = \underset{x}{\operatorname{argmax}}(\text{EIGF}(x)).$$

This point is selected as the new sample point for model evaluation and retraining of the GP surrogate model in the active learning loop shown in Fig. 1.

A closer look at the EIGF in Eq. (32) demonstrates a balance of two factors — often referred to as *exploration* and *exploitation*. In the first term, $(\hat{y}(x) - y(x^{(j^*)}))^2$, we see that the EIGF targets regions where variability in the function is large, which corresponds to exploitation. This is balanced by the second term, $\sigma_y^2(x)$, which corresponds to regions where the GP prediction has high uncertainty, likely caused by a lack of training data in this region and encouraging exploration. However, it is well-known that the EIGF learning function favors exploitation (i.e. the first term often dominates) and may not explore new regions of

parameter space sufficiently. Therefore a more balanced approach may be considered.

4.2. Variance of improvement for global fit (VIGF)

Mohammadi and Challenor [46] proposed a learning function using same improvement function as EIGF function, given in Eq. (31). They further recognized that $I(x)/\sigma^2(x)$ can be characterized by a non-central chi-square distribution [46].

$$I(x)/\sigma^2(x) \sim \chi^2 \left(\kappa = 1, \lambda = \left(\frac{\mu(x) - y(x^{(j^*)})}{\sigma(x)} \right)^2 \right) \quad (33)$$

where the number of degree of freedom (κ) and noncentrality parameter (λ) are defined using the GP's posterior mean and standard deviation. Thus, the variance of the improvement function is expressed as:

$$\text{VIGF}(x) = \text{Var}\{I(x)\} = 4\sigma^2(x) \left[(\mu(x) - y(x^{(j^*)}))^2 + 2\sigma^2(x) \right]. \quad (34)$$

In the VIGF, the first term $(4\sigma^2(x)(\mu(x) - y(x^{(j^*)}))^2)$ is a product of terms that contribute towards local exploitation and global exploration, thus favoring both. The second term $(8\sigma^4(x))$ focuses exclusively on global exploration. This learning function therefore gives more weight to global exploration than EIGF and may provide a better balance in the exploration/exploitation trade off.

4.3. MUSIC learning function

From Eqs. (26) and (30), we see that the quality of the Sobol index estimates depends strongly on the quality of the main effect GPs. Large uncertainties in the variance of the main effect GPs will result in associated large uncertainties in Sobol index estimates. We therefore aim to create a training set that produces the most accurate main effect GPs. This problem, as we will see, can be posed in active learning terms as identifying the set of points that yields the best global fit for the main effect GPs, taking inspiration from the EIGF [21] and VIGF [46].

Let us first consider that we can construct an improvement function for the main effect GPs in each dimension, similar to Eq. (31), as follows:

$$I_i(x_i) = (A_i(x_i) - E_{X_{-i}}[Y|X_i](x_i^{(j^*)}))^2, \quad i \in \{1, \dots, d\} \quad (35)$$

where x_i , $x_i^{(j^*)}$ are the i th component of points $\mathbf{x}$ and $\mathbf{x}^{(j^*)}$, respectively and $x_i^{(j^*)}$ is the nearest training sample to x_i in i th dimension. We immediately recognize that we cannot actually observe the true main effect $E_{X_{-i}}[Y(x_i^{(j^*)})|X_i]$. Instead, we must estimate this value using the main effect GP yielding a component-wise improvement function given by:

$$I_{A_i}(x_i) = (A_i(x_i) - A_i(x_i^{(j^*)}))^2, \quad i \in \{1, \dots, d\} \quad (36)$$

Taking the expectation and variance of Eq. (36) yields:

$$E[I_{A_i}(x_i)] = (\mu_{A_i}(x_i) - \mu_{A_i}(x_i^{(j^*)}))^2 + \sigma_{A_i}^2(x_i) \quad (37)$$

$$V[I_{A_i}(x_i)] = 4\sigma_{A_i}^2(x_i)[(\mu_{A_i}(x_i) - \mu_{A_i}(x_i^{(j^*)}))^2 + 2\sigma_{A_i}^2(x_i)] \quad (38)$$

where $\mu_{A_i}(x_i)$, $\mu_{A_i}(x_i^{(j^*)})$ can be computed using Eq. (22) and $\sigma_{A_i}^2(x_i)$ can be computed from the diagonal terms of Eq. (23). A closer look at $E[I_{A_i}(x_i)]$ demonstrates that it also balances exploration and exploitation specifically with regard to resolving the main effects. The first term targets regions where the main effects vary strongly with X_i (exploitation) and the second term targets regions where the main effect GP has high uncertainty due to a lack of training data (exploration). Eq. (38) has the same interpretation as the VIGF, balancing exploration and exploitation, but operating on the main effects.

Notice that the improvement criteria in Eqs. (37) and (38) focus on improving the prediction of individual main effect GPs. These learning criteria can be collectively represented by a vector as follows:

$$\mathbf{M}_A(\mathbf{x}) = \begin{bmatrix} E[I_{A_1}(x_1)] \\ E[I_{A_2}(x_2)] \\ \vdots \\ E[I_{A_d}(x_d)] \end{bmatrix} \quad \text{or} \quad \begin{bmatrix} V[I_{A_1}(x_1)] \\ V[I_{A_2}(x_2)] \\ \vdots \\ V[I_{A_d}(x_d)] \end{bmatrix} \quad (39)$$

From Eq. (39), we could directly derive improvement-based learning functions over individual main effects. This might be useful, for example, if we were interested in resolving a specific term in the HDMR in Eq. (1). However, we are not interested only in individual main effects. Rather, we are interested in ensuring convergence of all main effects, to the extent possible. This means we must somehow balance the d improvement functions that result from Eq. (39) to establish a single improvement function. We do so by introducing a distance-based pre-factor that combines individual learning criteria on main effect GPs. This results in a new learning function, expressed as:

$$\mathcal{J}(\mathbf{x}) = \mathbf{D}^T(\mathbf{x})\mathbf{M}_A(\mathbf{x}) \quad (40)$$

and termed the MUSIC (Minimizing Uncertainty in Sobol Index Convergence) learning function. We can consider numerous different distance pre-factors. Here, we specifically propose the following two pre-factors that produce convex combinations of main effect improvements:

$$\mathbf{D}_1(\mathbf{x}) = \mathbf{W} \cdot \begin{bmatrix} |x_1 - x_1^{(j^*)}| \\ |x_2 - x_2^{(j^*)}| \\ \vdots \\ |x_d - x_d^{(j^*)}| \end{bmatrix} \quad \text{and} \quad \mathbf{D}_2(\mathbf{x}) = \mathbf{W} \cdot \|\mathbf{x} - \mathbf{x}^{(j^*)}\|_2^2 \mathbf{1}_d \quad (41)$$

where $\|\mathbf{x} - \mathbf{x}^{(j^*)}\|_2$ is the Euclidean distance between the sample $\mathbf{x}$ and nearest point $\mathbf{x}^{(j^*)}$, and $|x_i - x_i^{(j^*)}|$ is the distance between the i th component of $\mathbf{x}$ and nearest point $\mathbf{x}^{(j^*)}$ along the i th dimension. In either function, given two points whose total main effect improvements are comparable the algorithm should favor the point that is furthest from its nearest neighbor either in a Euclidean sense or in a component-wise sense. The weights ($\mathbf{W} = [w_1, w_2, \dots, w_d]^T$) can be arbitrarily selected such that $\sum_i w_i = 1$ to favor improvement in particular main

Table 1

Four different combinations of MUSIC learning strategy based on improvement and distance-based pre-factor.

S.No.	Combinations	Learning function
1	MUSIC+EIGF D1	$\mathcal{J}_{E_1}(\mathbf{x}) = \sum_{i=1}^d w_i x_i - x_i^{(j^*)} E[I_{A_i}(x_i)]$
2	MUSIC+EIGF D2	$\mathcal{J}_{E_2}(\mathbf{x}) = \ \mathbf{x} - \mathbf{x}^{(j^*)}\ _2^2 \sum_{i=1}^d w_i E[I_{A_i}(x_i)]$
3	MUSIC+VIGF D1	$\mathcal{J}_{V_1}(\mathbf{x}) = \sum_{i=1}^d w_i x_i - x_i^{(j^*)} V[I_{A_i}(x_i)]$
4	MUSIC+VIGF D2	$\mathcal{J}_{V_2}(\mathbf{x}) = \ \mathbf{x} - \mathbf{x}^{(j^*)}\ _2^2 \sum_{i=1}^d w_i V[I_{A_i}(x_i)]$

effects. For example, these can be set equal to the existing Sobol index estimates in order to favor exploration in directions that correspond to higher Sobol indices.

Using this learning function, we then select the next sample as the one that maximizes it as:

$$\mathbf{x}^* = \underset{\mathbf{x}}{\operatorname{argmax}} \mathcal{J}(\mathbf{x})$$

from among the set of n_j candidate samples. Table 1 summarizes the learning function for four combinations based on the improvement criteria and the distance-based pre-factor.

5. Exploration of convergence using analytical functions

In this section, we study the convergence of the proposed active learning methods for four benchmark analytical functions with known Sobol indices and having increasing dimensions ranging from 2 to 15. Prior to presenting these examples, we provide a simple basis upon which to explain the observed convergence by considering the ratio of two quantities with some error.

5.1. Ratios of errors

The dependence on the absolute error of a ratio on estimates of the numerator and denominator is not straightforward. As we will see, reducing error in the numerator, denominator, or both does not necessarily reduce error in the ratio. Consider the following ratio:

$$S = \frac{A}{Y}$$

Next, consider that we have estimates of S , A and Y given by $\hat{S}$, $\hat{A}$ and $\hat{Y}$. Define the absolute error in the numerator (A) and denominator (Y) as:

$$\Delta A = |A - \hat{A}|$$

$$\Delta Y = |Y - \hat{Y}|$$

such that

$$\hat{A} = A \pm \Delta A$$

$$\hat{Y} = Y \pm \Delta Y$$

Next, the absolute error in the ratio can be expressed as:

$$\Delta S = |S - \hat{S}| = \left| S - \frac{\hat{A}}{\hat{Y}} \right|$$

$$\Delta S = \left| S - \frac{A \pm \Delta A}{Y \pm \Delta Y} \right|$$

$$\Delta S = \left| \frac{S(Y \pm \Delta Y) - A \mp \Delta A}{Y \pm \Delta Y} \right|$$

$$\Delta S = \left| \frac{SY \pm S\Delta Y - A \mp \Delta A}{Y \pm \Delta Y} \right|$$

$$\Delta S = \left| \frac{\pm S\Delta Y \mp \Delta A}{Y \pm \Delta Y} \right|,$$

which depends, in a nontrivial way, on the errors ΔA and ΔY – but also depends on the value of S itself.

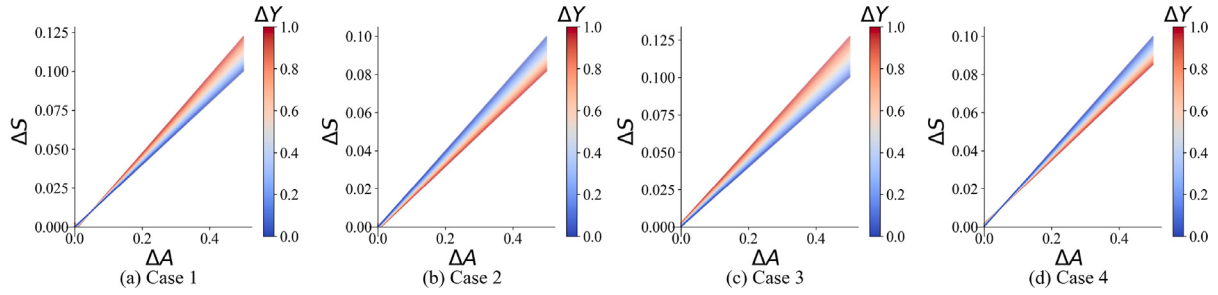


Fig. 2. Influence of errors in the numerator, ΔA , and denominator, ΔY , on error ΔS in the ratio $S = A/Y$ for $S = 0.01$ for four cases of systematic bias (a) Case 1: Both ΔA and ΔY are underestimated, (b) Case 2: Both ΔA and ΔY are overestimated, (c) ΔA is overestimated and ΔY is underestimated, and (d) ΔA is underestimated and ΔY is overestimated.

Recognizing that our GP models for A and Y presented in the formulation of the GP-based Sobol indices above may be systematically biased (depending on the training set and learning function being used) and that we may not know this bias a priori, we explore this error under four specific cases.

- Case 1: $\hat{A}$ and $\hat{Y}$ are both underestimated such that $\hat{A} = A - \Delta A$ and $\hat{Y} = Y - \Delta Y$ yielding

$$\Delta S = \frac{|-S\Delta Y + \Delta A|}{Y - \Delta Y}$$

- Case 2: $\hat{A}$ and $\hat{Y}$ are both overestimated such that $\hat{A} = A + \Delta A$ and $\hat{Y} = Y + \Delta Y$ yielding

$$\Delta S = \frac{|S\Delta Y - \Delta A|}{Y + \Delta Y}$$

- Case 3: $\hat{A}$ is overestimated ($\hat{A} = A + \Delta A$) and $\hat{Y}$ is underestimated ($\hat{Y} = Y - \Delta Y$) yielding

$$\Delta S = \frac{|-S\Delta Y - \Delta A|}{Y - \Delta Y}$$

- Case 4: $\hat{A}$ is underestimated ($\hat{A} = A - \Delta A$) and $\hat{Y}$ is overestimated ($\hat{Y} = Y + \Delta Y$) yielding

$$\Delta S = \frac{|S\Delta Y + \Delta A|}{Y + \Delta Y}$$

Figs. 2 and 3 show plots of the dependence of ΔS on ΔA and ΔY for $S = 0.01$ and $S = 0.8$, respectively, in each of these four cases for a hypothetical ratio. Notice that, for $S = 0.01$ (Fig. 2) the errors appear to decrease towards zero as both ΔA and ΔY decrease.¹ That is, if either (or both) ΔA and ΔY decrease, then ΔS decreases. However, for $S = 0.8$, a decrease in either ΔA or ΔY (or both) does not guarantee a decrease in ΔS . In Cases 1 and 2, regardless of ΔA (even for very high ΔA), it is possible to obtain zero ΔS . In fact, for a fixed value of ΔY , as ΔA decreases ΔS will decrease to zero and then begin increasing again. This occurs because the systematic bias in both Case 1 and Case 2 can cause the ratio $S = A/Y$ to be correct even if both A and Y are (perhaps severely) incorrect. Indeed, two points X_1 and X_2 are shown in Fig. 3 where $\Delta A_{X_1} > \Delta A_{X_2}$ and $\Delta Y_{X_1} > \Delta Y_{X_2}$, yet $\Delta S_{X_1} < \Delta S_{X_2}$. As we will see, this systematic bias in A and Y can make accelerated convergence in the Sobol indices difficult.

Next, we will study the convergence, and specifically the behavior of errors in light of the discussion above, for the different active learning schemes using four analytical test cases of increasing dimension.

5.2. Square exponential function

First, we consider the following square exponential function of two random inputs (i.e. $\mathbf{X} = [X_1, X_2]$).

$$y = X_1 \exp(-X_1^2 - X_2^2)$$

¹ In fact, Cases 1 and 2 have the same bias observed in the $S = 0.8$ case but the influence is negligible because S is very small.

where, X_1 and X_2 are uniformly distributed between $[a, b]$. Two cases are considered in this study, one with tight bounds around the two peaks ($a = -2$ and $b = 2$) and another with larger bounds ($a = -2$ and $b = 6$) such that a large flat region is observed across much of the domain. Both cases are illustrated in Fig. 4. The presence of this flat region has a significant impact on the Sobol indices and, moreover, affects the convergence of the active learning algorithms. In both cases, the Sobol indices can be computed analytically. For the first case ($b = 2$), the first-order Sobol indices are $S_1 = 0.6208$ and $S_2 = 0$. For the second case, the Sobol indices are $S_1 = 0.3119$ and $S_2 = 2.30 \times 10^{-5}$.

Fig. 5 shows the mean square error convergence from 100 repeated trials of the different active learning schemes using the EIGF and MUSIC learning functions compared with random sampling for the case where $b = 2$ (no flat region). Fig. 5(a) shows convergence of the variance of the output, σ_Y^2 , as estimated by the GP. We see that the EIGF generally outperforms the other two approaches, converging to a very accurate estimate for σ_Y^2 with less than ~ 30 samples. Fig. 5(b) and (c) show the same convergence plots for the variance of the main effect GPs $\sigma_{A_1}^2$ and $\sigma_{A_2}^2$, respectively. We see that, in main effects the MUSIC learning function offers comparable convergence to the EIGF in A_1 and superior convergence in A_2 . Both methods generally outperform random sampling in the main effect variances.

Finally, Fig. 5(d) and (e) show convergence of the Sobol indices S_1 and S_2 . For this problem, the Sobol convergence tracks closely with the main effect convergence, but with slight differences. For example, we recall that the variance $\sigma_{A_1}^2$ from Fig. 5(b) is the numerator and σ_Y^2 from Fig. 5(a) is the denominator of the Sobol index S_1 . However, despite superior performance in both $\sigma_{A_1}^2$ and σ_Y^2 , the MUSIC function performs mildly worse than random sampling in S_1 for small samples. This is due to the ‘ratio of errors’ issues discussed in the previous section.

Fig. 6 shows the same convergence plots but compares VIGF and MUSIC (based on VIGF) learning functions against the sequential random sampling. Similar to EIGF, the VIGF learning function results in the best estimate for output variance (σ_Y^2), as shown in Fig. 6(a). However, VIGF outperforms EIGF for main effect and sobol estimates ($\sigma_{A_i}^2$) and yields similar convergence to the ‘MUSIC+VIGF D2’ combination. This improvement is due to the better prediction at both functional features (peak and valley). Consider, the main effect of first dimension $E[Y|X_1]$ for a fixed value of X_1 (imagine a cross-section at arbitrary X_1 in Fig. 4), the main effect function depends on either the peak or the valley (depending on the value of X_1). But, the main effect of the second dimension ($E[Y|X_2]$) depends on both features for all values of X_2 . Since VIGF gives more importance to exploration (as compared to EIGF), it resolves both features simultaneously — it does not get stuck exploiting one feature. This results in better estimates for σ_Y^2 as well as $\sigma_{A_1}^2$ and eventually better Sobol estimates as shown in 6(e).

Figs. 7 and 8 show the same convergence plots for the case with $b = 6$ (having a flat region). As expected, the EIGF/VIGF learning functions perform much better than the other methods for σ_Y^2 , as they both perform well for tasks where exploitation is paramount [50]. They both also exhibit superior performance in the main effect variances $\sigma_{A_1}^2$

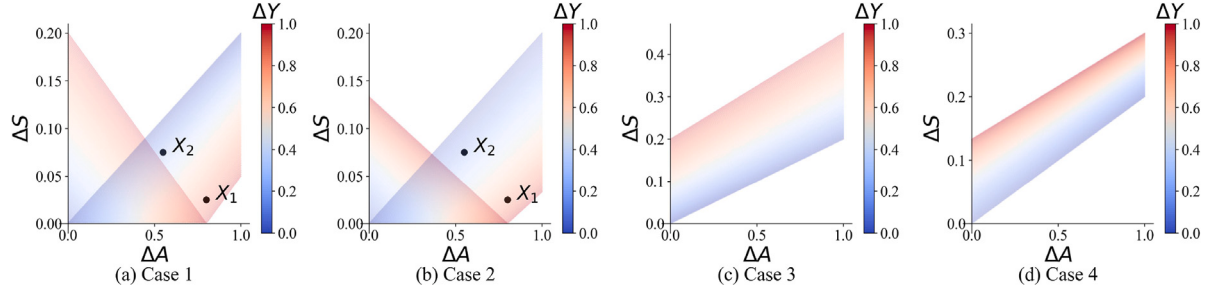


Fig. 3. Influence of errors in the numerator, ΔA , and denominator, ΔY , on error ΔS in the ratio $S = A/Y$ for $S = 0.8$ for four cases of systematic bias (a) Case 1: Both ΔA and ΔY are underestimated, (b) Case 2: Both ΔA and ΔY are overestimated, (c) ΔA is overestimated and ΔY is underestimated, and (d) ΔA is underestimated and ΔY is overestimated.

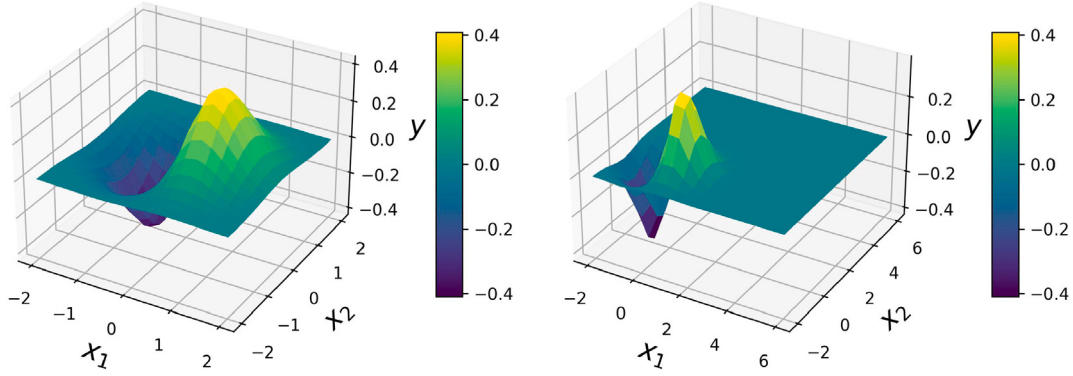


Fig. 4. Square exponential functions (a) with tight bounds and no flat region and (b) with wide bounds and large flat region.

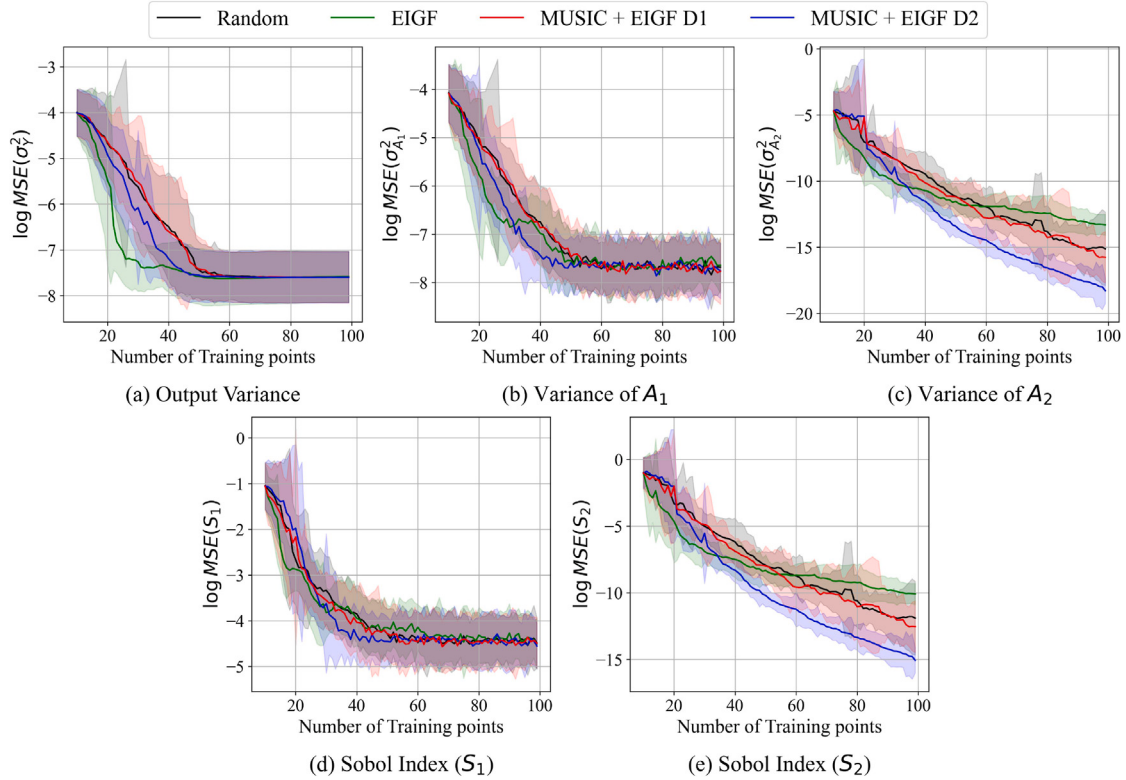


Fig. 5. Exponential Function ($b = 2$, no flat region) – Mean square convergence with σ confidence intervals from 100 repeated trials of (a) output variance σ_y^2 , (b,c) main effect variances $\sigma_{A_1}^2$ and $\sigma_{A_2}^2$, and (d,e) Sobol indices S_1 and S_2 for MUSIC and EIGF active learning schemes compared with random sampling.

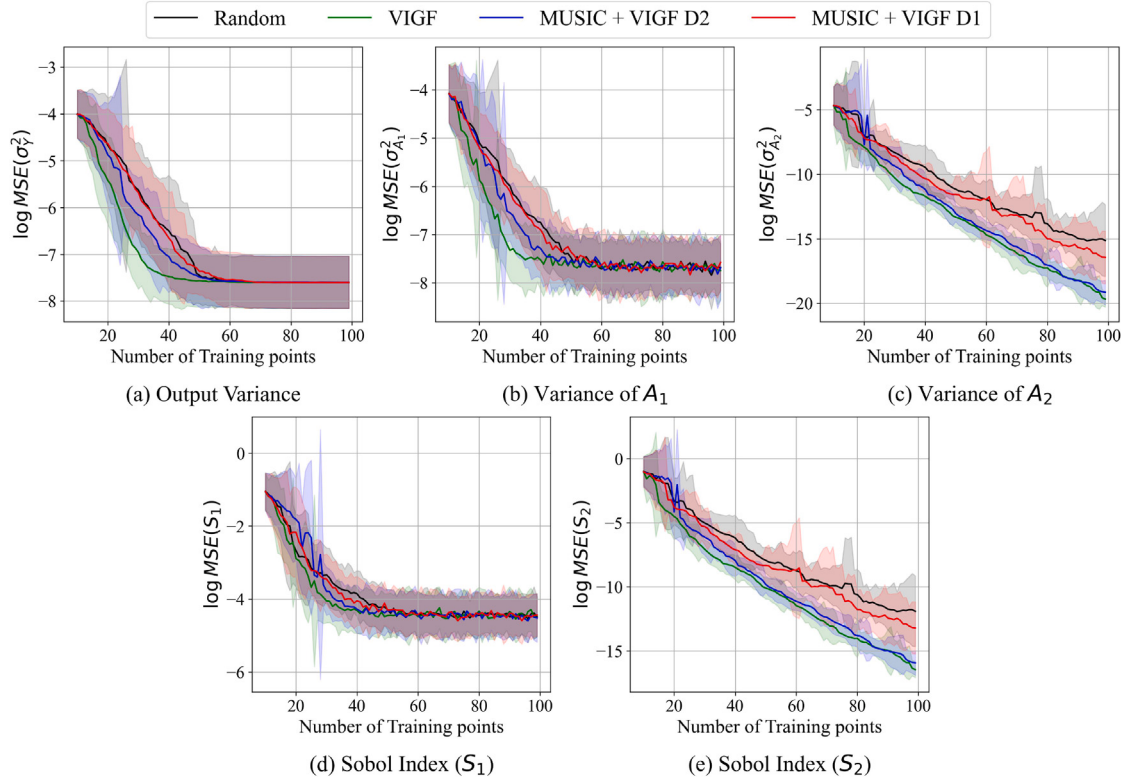


Fig. 6. Exponential Function ($b = 2$, no flat region) – Mean square convergence with σ confidence intervals from 100 repeated trials of (a) output variance σ_y^2 , (b,c) main effect variances $\sigma_{A_1}^2$ and $\sigma_{A_2}^2$, and (d,e) Sobol indices S_1 and S_2 for MUSIC and VIGF active learning schemes compared with random sampling.

and $\sigma_{A_2}^2$ and Sobol indices S_1 and S_2 for small sample sizes. Again, this is a result of its strong exploitation properties. The MUSIC function meanwhile provides superior convergence to random sampling in all cases as well.

5.3. Ishigami function

Next, consider the well-known Ishigami function with three independent input random variables ($\mathbf{X} = [X_1, X_2, X_3]$) that are uniformly distributed over $[-\pi, \pi]$, and defined by:

$$y(\mathbf{x}) = \sin x_1 + a \sin^2 x_2 + b x_3^4 \sin x_1$$

where parameters $a = 7$, and $b = 0.1$. This function has nonlinear and nonmonotonic behavior, but the third input dimension exhibits a peculiar behavior. The first-order contribution of the third dimension is zero towards output variance, but it cannot be completely ignored as an interaction sensitivity index (i.e. S_{13}) has a significant impact on output variance. The first-order Sobol indices are $S_1 = 0.3139$, $S_2 = 0.4424$, and $S_3 = 0$, determined analytically by solving for the variance of the main effect and output as:

$$\text{Var}\{A_1\} = \frac{1}{2} \left(1 + \frac{b\pi^4}{5} \right)^2, \quad \text{Var}\{A_2\} = \frac{a^2}{8}, \quad \text{Var}\{A_3\} = 0$$

$$\text{Var}\{Y\} = \frac{a^2}{8} + \frac{b\pi^4}{5} + \frac{b^2\pi^8}{18} + \frac{1}{2}$$

The performance of EIGF and MUSIC ('EIGF D1' and 'EIGF D2') adaptive sampling strategies is shown in Fig. 9. The adaptive algorithms start with a small set of 10 samples generated using Latin Hypercube Sampling and convergence is observed up to 500 samples. After updating the surrogate model, Sobol indices are estimated using 25,000 candidate points. The process is repeated for 100 trials to obtain confidence intervals in the convergence. In this example, the EIGF struggles to accurately estimate the output variance (σ_y^2 in Fig. 9(a)) and the

main effect variances ($\sigma_{A_i}^2$ in Fig. 9(b)–(d)). This can be explained by the sinusoidal nature of the Ishigami function, which requires a balance between exploration and exploitation. Error convergence of the 'MUSIC+EIGF D1' learning function is similar to random sampling for every estimate. 'MUSIC+EIGF D2', on the other hand, leads to superior convergence in nearly all aspects. The distance-based pre-factor is the main difference here. The D_2 function uses the Euclidean distance and ensures that training samples are globally far away from each other, which helps the 'MUSIC+EIGF' learning function to avoid local exploitation. Because the D_1 function utilizes an element-wise distance component, it does not avoid local exploitation as effectively. Fig. 9(e–g) show the ultimate convergence of the Sobol index estimates where again 'MUSIC+EIGF D2' generally outperforms the other criteria.

Fig. 10 likewise compares the VIGF and MUSIC (+VIGF) learning functions to random sampling. The main difference here is that the VIGF improves performance considerably over the EIGF. Given this consistent improvement with the VIGF, we will compare only the results using VIGF in the remaining analytical examples.

Here, the convergence of the first and second Sobol indices may appear counterintuitive. The MUSIC function estimates both the total and main effect variances with high accuracy, but the error in the Sobol estimates is comparable to random sampling. This is again due to the effect of ratios of errors. Fig. 11 shows the convergence path of the Sobol indices by plotting traces of $(\hat{\sigma}_Y^2, \hat{\sigma}_{A_i}^2)$ averaged over 100 trials. The gray surface represents the Sobol index for any specific pair. The green star is the actual first-order Sobol index (S_i) value for the i^{th} dimension. The gray line shows all combinations of $(\hat{\sigma}_Y^2, \hat{\sigma}_{A_i}^2)$, such that their ratio is equal to the true Sobol index (i.e. $\hat{S}_i = S_i$). The black, orange, and blue curves illustrate the convergence of random, EIGF, and MUSIC sampling, respectively. We notice that, although MUSIC and EIGF perform better in the variance estimators (see convergence above), the path that their estimators takes towards the true value does

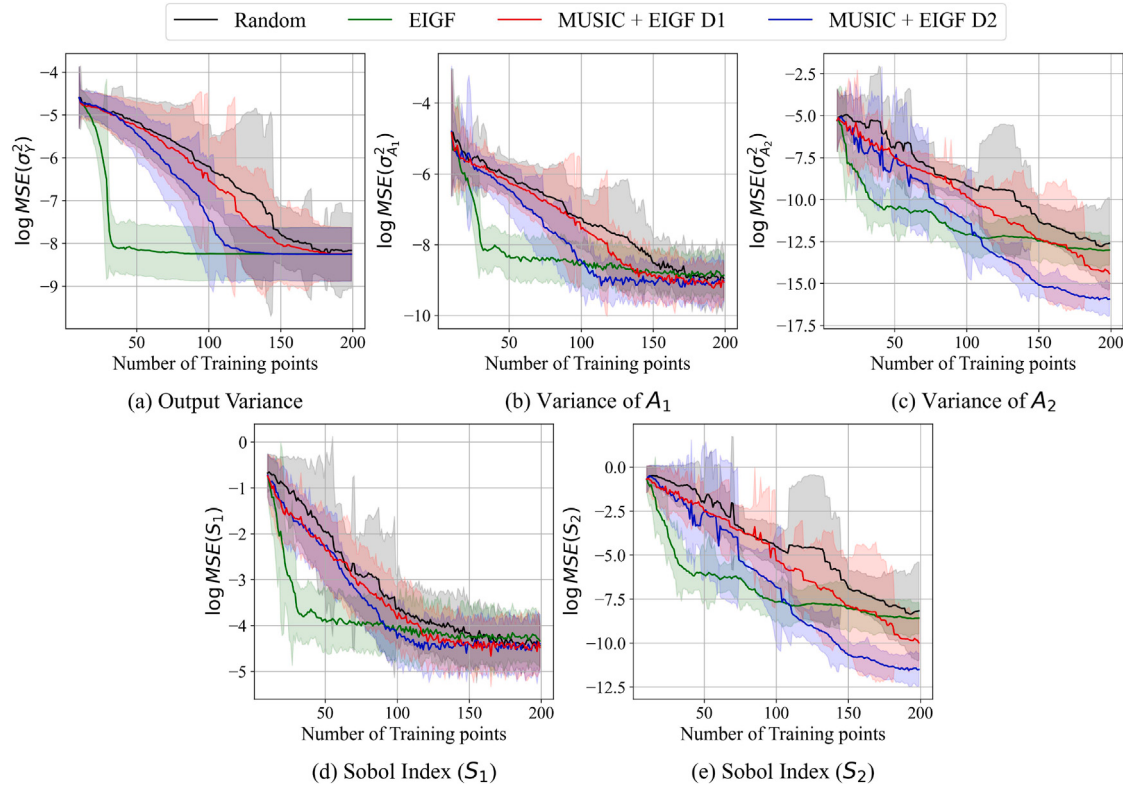


Fig. 7. Exponential Function ($b = 6$, with flat region) – Mean square convergence with σ confidence intervals from 100 repeated trials of (a) output variance σ_y^2 , (b,c) main effect variances $\sigma_{A_1}^2$ and $\sigma_{A_2}^2$, and (d,e) Sobol indices S_1 and S_2 for MUSIC and EIGF active learning schemes compared with random sampling.

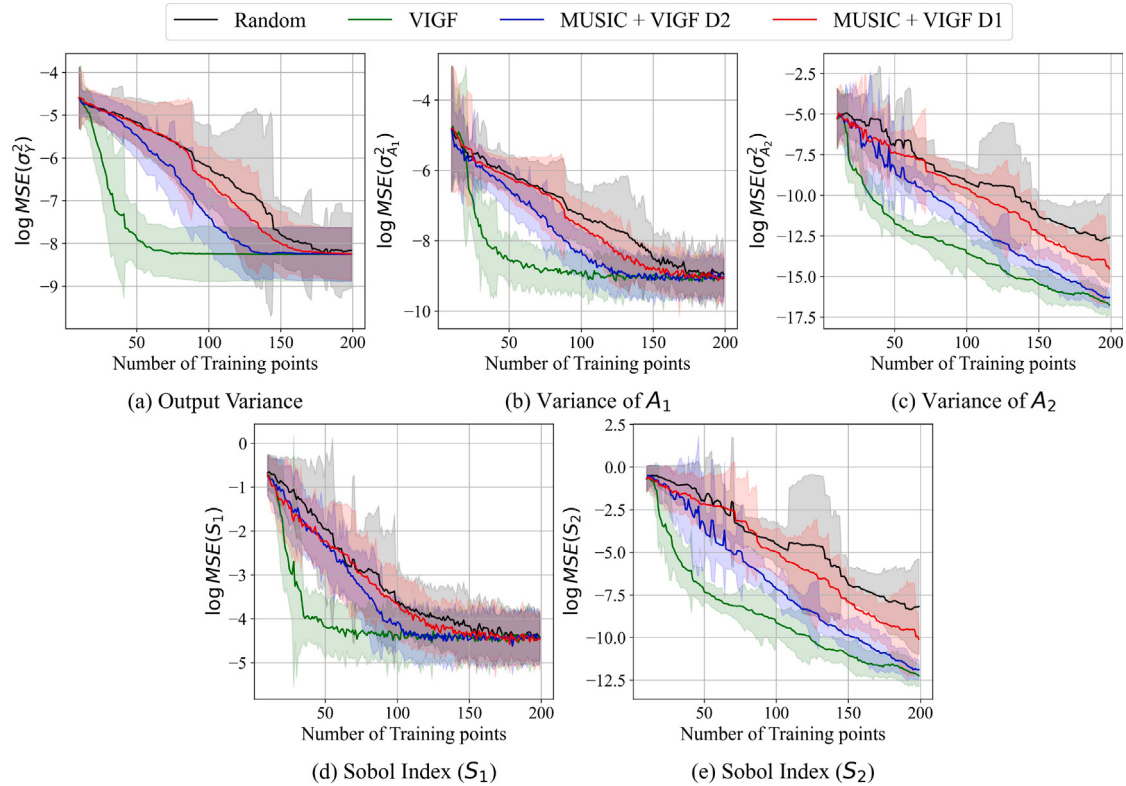


Fig. 8. Exponential Function ($b = 6$, with flat region) – Mean square convergence with σ confidence intervals from 100 repeated trials of (a) output variance σ_y^2 , (b,c) main effect variances $\sigma_{A_1}^2$ and $\sigma_{A_2}^2$, and (d,e) Sobol indices S_1 and S_2 for MUSIC and VIGF active learning schemes compared with random sampling.

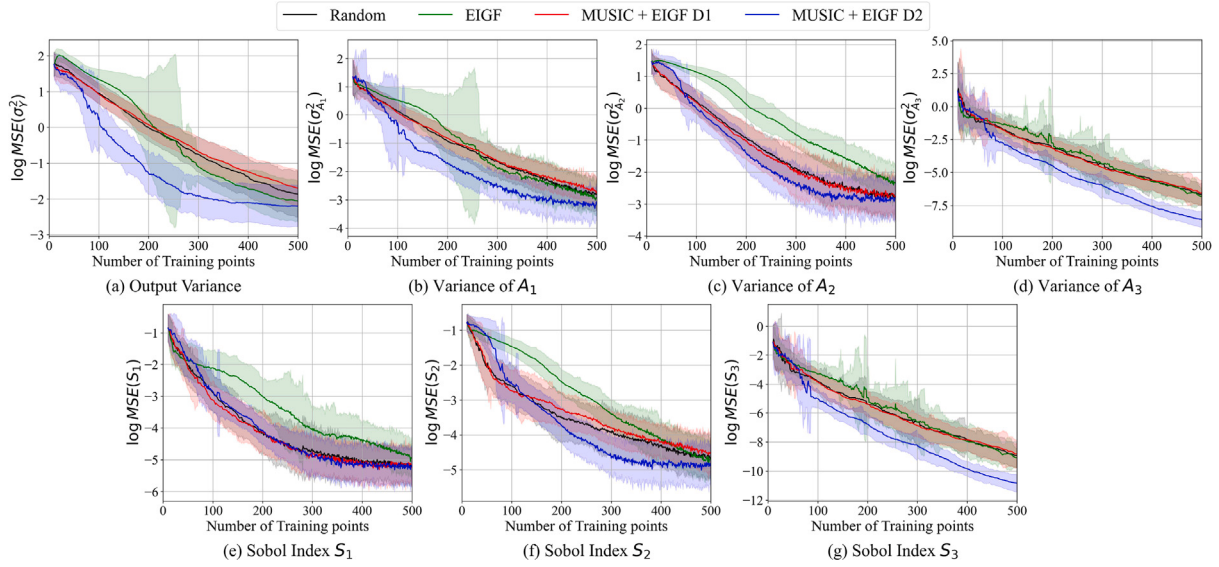


Fig. 9. Ishigami Function — Mean square convergence with σ confidence intervals from 100 repeated trials of (a) output variance σ_y^2 , (b–d) main effect variances $\sigma_{A_i}^2 \forall i \in \{1, 2, 3\}$, and (e–g) Sobol indices S_1 , S_2 , and S_3 for MUSIC and EIGF active learning schemes compared with random sampling.

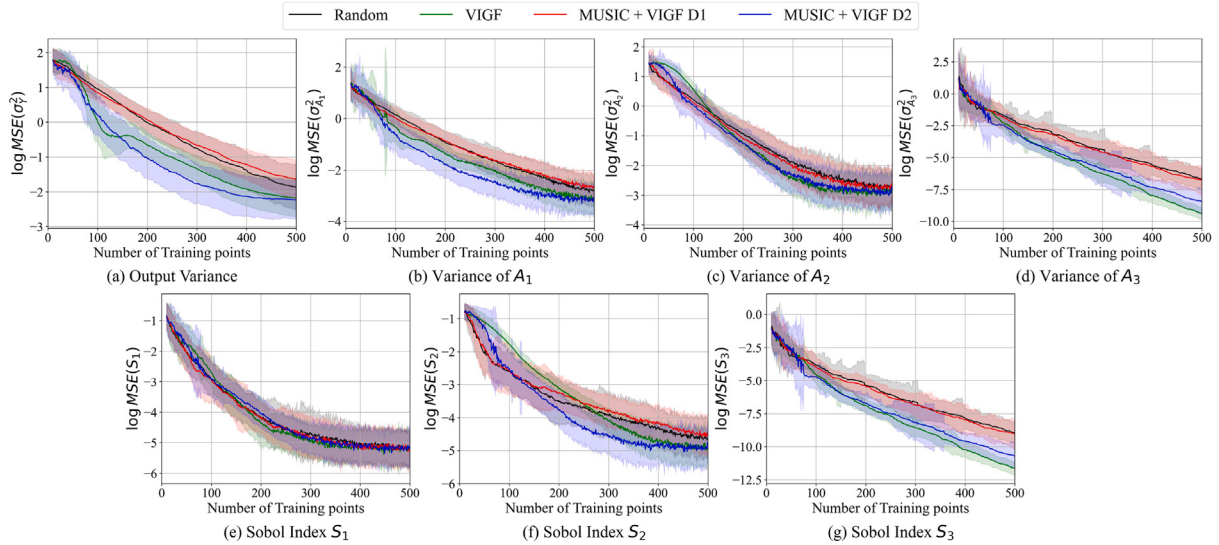


Fig. 10. Ishigami Function — Mean square convergence of (a) output variance σ_y^2 , (b–d) main effect variances $\sigma_{A_i}^2 \forall i \in \{1, 2, 3\}$, and (e–g) Sobol indices S_1 , S_2 , and S_3 for MUSIC and VIGF active learning schemes compared with random sampling.

not follow the gray line. That is, their Sobol index estimates are biased due to the ratio of errors. This is not the case for random sampling, which converges largely along the gray line. This counter-intuitive behavior is not observed for the third dimension because the actual Sobol index is zero (i.e. $S_3 = 0$) and therefore depends only on the main effect GP. As illustrated in Fig. 2, when S_i is very small convergence in the numerator will yield convergence in the ratio. We can therefore conclude that the MUSIC function will be very effective at identifying small Sobol indices.

5.4. G-function

Next, we consider the 5-dimensional G-function, defined as the following product of non-linear functions in each dimension,

$$y(\mathbf{x}) = \prod_{k=1}^d \frac{|4x^{(k)} - 2| + a^{(k)}}{a^{(k)} + 1}$$

where each input feature is uniformly distributed between $[0, 1]$ (i.e. $X^{(k)} \sim \text{Uniform}(0, 1)$). The analytical Sobol indices in five dimensions (i.e. $d = 5$) are 0.48, 0.21, 0.12, 0.08, 0.06, corresponding to $a^{(k)} = k$, $k \in 1, 2, 3, 4, 5$ from the following equations:

$$\begin{aligned} \text{Var}\{A_i\} &= \frac{1}{3(1 + a^{(i)})^2} & \text{Var}\{Y\} &= \left[\prod_{k=1}^d (1 + \text{Var}\{A_k\}) \right] - 1 \\ S_i &= \frac{1/(3(1 + a^{(i)})^2)}{\left[\prod_{k=1}^d [1 + 1/(3(1 + a^{(k)})^2)] \right] - 1} \end{aligned}$$

Again we compare the three sampling strategies, where each sampling strategy initiates by training a GP on 30 samples obtained from Latin Hypercube Sampling and adaptively selects new samples up to 500 samples. The Sobol indices are computed using 25 000 candidate points. Fig. 12 shows the convergence for VIGF and MUSIC (based on VIGF) compared to random sampling. Here, the ‘MUSIC+VIGF’ strategy and the VIGF show superior performance to random sampling in the total variance (Fig. 12(a)), main effects (Fig. 12(b)–(f)), and Sobol

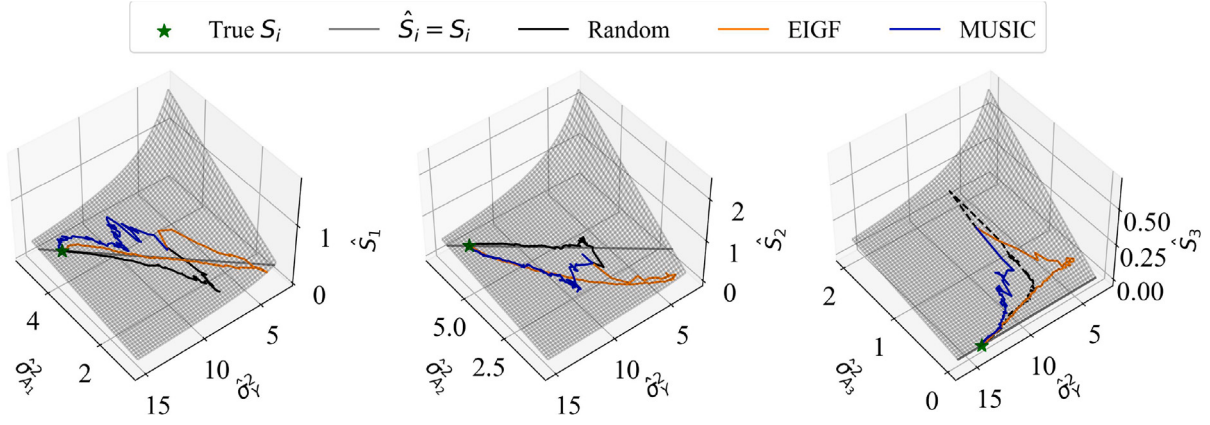


Fig. 11. Dependence of Sobol indices on the variance of main effect GPs and Output GP.

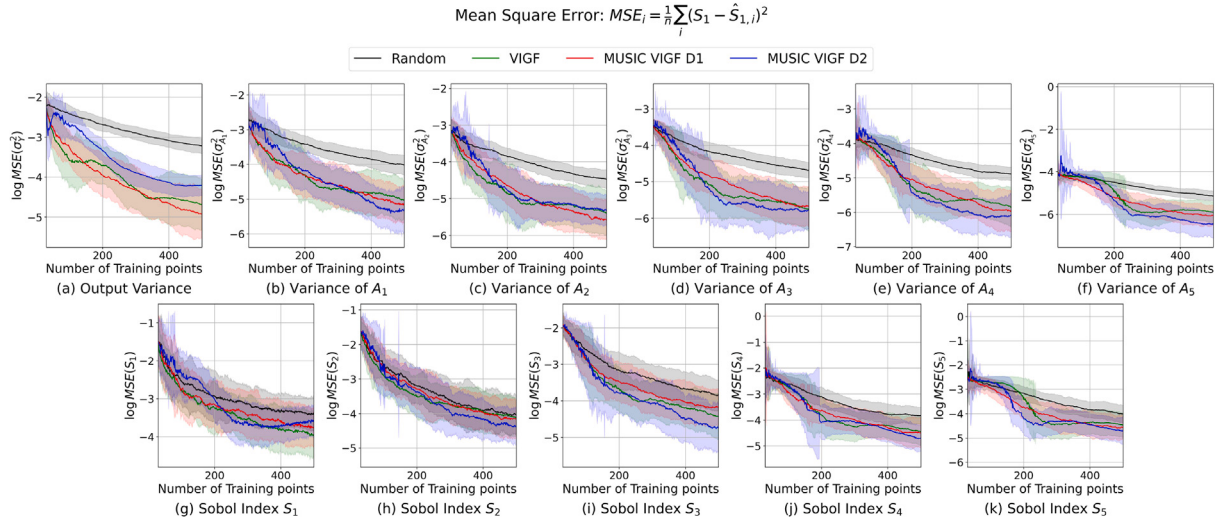


Fig. 12. G-Sobol Function — Mean square convergence with σ confidence intervals from 100 repeated trials of (a) output variance σ_y^2 , (b–f) main effect variances $\sigma_{A_i}^2 \forall i \in \{1, 2, 3, 4, 5\}$, and (g–k) Sobol indices S_i for MUSIC and VIGF active learning schemes compared with random sampling.

indices ((Fig. 12(g)–(k)). Results from EIGF learning are not shown due to poor convergence in comparison to VIGF.

5.5. Gaussian function

To conclude, we compare performance for a high input dimensional problem. A synthetic 15d function is considered here, such that 75.5% of the contribution is from the first-order effects of three inputs. The function is defined as a product of 15 independent square exponential functions, similar to a Gaussian kernel given by:

$$y(\mathbf{x}) = \prod_{i=1}^d \exp\left(-\frac{x_i^2}{a_i}\right)$$

where, $x_i \sim U(-3, 3) \forall i \in \{1, 2, \dots, d\}$ and

$$\mathbf{a} = [1.45, 3.3, 15, 50, 55, 58, 59, 100, 102, 112.5, 150, 160, 180, 190, 200]$$

Fig. 13(a) shows the main effect contribution and Fig. 13(b) shows the true Sobol indices (in log scale).

Fig. 14 shows convergence of the Sobol indices for the first three dimensions using the VIGF and MUSIC with VIGF learning functions compared to random sampling for 500 samples. We notice that the VIGF and ‘MUSIC+VIGF D2’ perform quite poorly for this example. In particular, the ‘MUSIC+VIGF D2’ exhibits poor convergence because the Euclidean distance is not a good measure of distance in 15 dimensions [51]. When this is replaced with the component-wise

distance measure, convergence is much better and is comparable with random sampling. Convergence for dimensions 4–15 are not shown here because they immediately produce very low errors in all cases given that the corresponding Sobol indices are near zero. Importantly, we observe that no learning function is capable of producing superior convergence to random sampling in high-dimensions.

6. Application to boundary layer wind tunnel experimental design

The original motivation for developing this active learning scheme for Sobol index estimation came from an experimental investigation into the influence of roughness terrain parameters in large-scale Boundary Layer Wind Tunnel (BLWT) experiments. In this study, experiments were performed using the University of Florida BLWT (UF-BLWT illustrated in Fig. 15), which features a novel automated roughness element grid called the Terraformer. This Terraformer is capable of rapidly and automatically changing the roughness terrain at the push of a button to mimic a wide-range of wind flow conditions [52].

In this work, we parameterize the Terraformer (which has 1116 independent elements capable of changing both height and profile orientation) by a random field in the along-wind direction using a truncated Karhunen–Loeve (KL) expansion such that the element at a location x along the Terraformer is given by:

$$h(x, \theta) = \mu_H + \sum_{i=1}^d \sqrt{\lambda^{(i)}} \theta^{(i)} f^{(i)}(x)$$

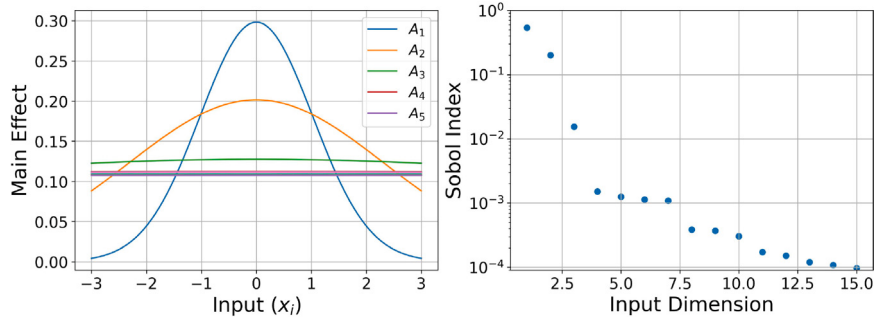


Fig. 13. Gaussian function: (a) First-order main effect functions and (b) True Sobol Indices.

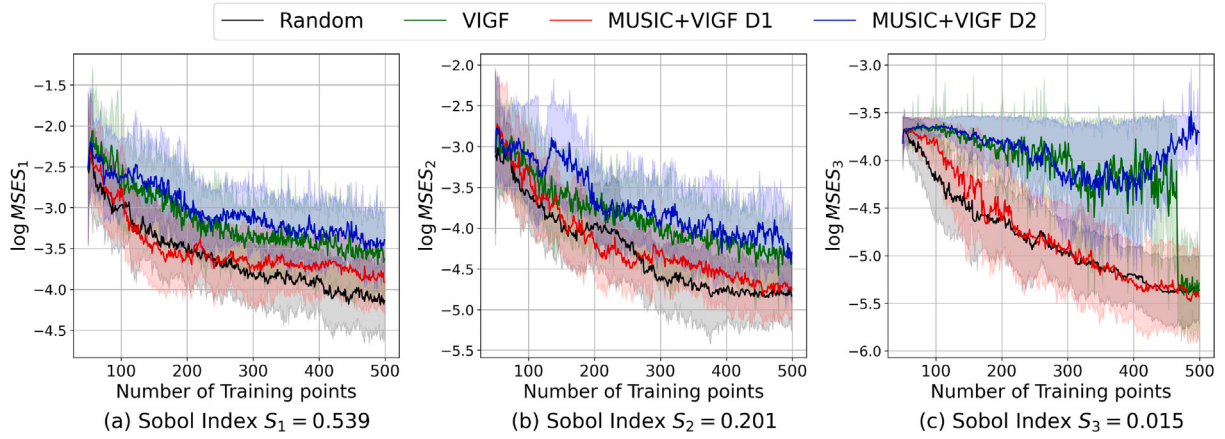


Fig. 14. Gaussian Function — Mean square convergence with σ confidence intervals from 50 repeated trials for Sobol indices (a) S_1 , (b) S_2 , and (c) S_3 for MUSIC and VIGF active learning schemes compared with random sampling.

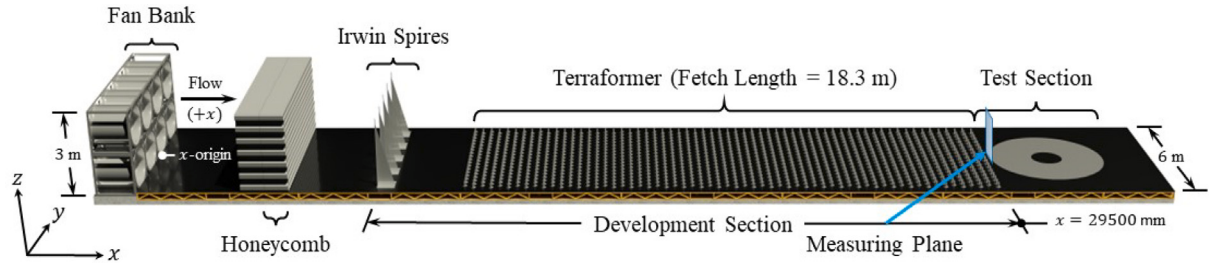


Fig. 15. Schematic of the University of Florida Boundary Layer Wind Tunnel (BLWT). The Terraformer roughness element grid is capable of rapidly and automatically reconfiguring its 1116 elements to produce a wide-range of wind flow characteristics.

where $\mu_H = 80$ mm is the mean element height, $\theta^{(i)}$ are standard normal random variables, and $\lambda^{(i)}$ and $f^{(i)}(x)$ are the eigenvalues and eigenfunctions of the covariance function given by:

$$C(x_1, x_2) = \sigma^2 \exp(-a|x_2 - x_1|) \cos(\omega|x_2 - x_1|) \quad (42)$$

Here $\sigma^2 = 100$ mm is the variance of the field, ω is the wave number, and a and ω are selected such that the length scale of the covariance function is given by $L = a/(a^2 + \omega^2) = 3000$ mm, or approximately 1/6 of the Terraformer length. The KL expansion is truncated to $d = 10$ terms.

The significant expense in this project is due to experimental time, where the execution time for a single experiment is approximately 20 min. Even with the high degree of automation afforded by the UF-BLWT, we cannot afford to run a sufficient number of experiments to explore the influence of all $d = 10$ random field parameters on the resulting wind flow. We therefore aim to adaptively perform global sensitivity analysis to identify which random field parameters have the most influence on the wind characteristic of interest and reduce the

dimension of the random field to facilitate a feasible set of experiments. Our specific characteristic of interest is the turbulence intensity profile of the resulting wind flow defined at a height above the floor (z) by:

$$I_u(z) = \frac{\sigma_u(z)}{\mu_u(z)} \quad (43)$$

where $\mu_u(z)$ and $\sigma_u(z)$ are the mean and standard deviation of the longitudinal component of the wind velocity $u(z)$ – which is a random process. More specifically, we aim to quantify the sensitivity of the following distance measure to the 10 KL random variables:

$$d(\theta) = \|I_u(\theta) - I_u^*\|_2 \quad (44)$$

where $I_u(\theta)$ is the vector of discretized turbulence intensities observed at points separated by 20 mm in the range $z = [180, 500]$ mm for a Terraformer configuration with parameters θ and I_u^* is the same vector defined for a reference roughness grid with all element uniformly set to $h = 80$ mm. In other words, we aim to quantify the sensitivity of the

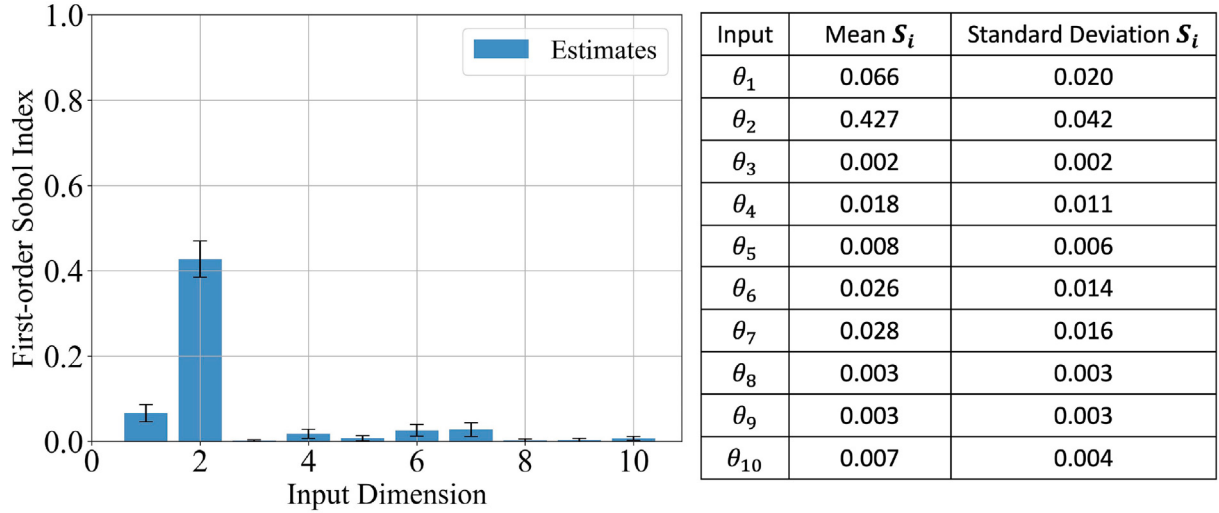


Fig. 16. Final estimates of the first-order Sobol Indices.

deviation of the turbulence intensity profile from the reference profile to each of the KL random variables.

For this investigation, we applied an earlier component-wise application of the proposed EIGF-based MUSIC learning function not demonstrated in detail above. This initial version of MUSIC focuses on only one dimension at a time and chooses a new sample to minimize the uncertainty in the estimate of the particular Sobol index associated with the largest MUSIC component. More specifically, using the expected main effect improvement defined in Eq. (37), this version selects the dimension with maximum improvement in the individual main effect GPs as:

$$x_i^* = \operatorname{argmax}_{x_i} E[I_{A_i}(x_i)] \quad (45)$$

where x_i^* is the i th element of the new sample $\mathbf{x}^*$ and the other elements are selected randomly over the input marginal distribution. Although this version of the MUSIC function was ultimately improved and the methods discussed in detail above were explored in more detail, it was successfully used to adaptively estimate sensitivities for the BLWT experimental investigation.

Initially, 30 experiments were performed with parameters θ selected by Latin Hypercube sampling. Then, 270 more experiments were conducted using the MUSIC learning strategy described above. Fig. 16 shows the final estimates (after 300 experiments) of the first-order Sobol indices, along with the standard error in the estimates. The experiments show that the distance metric between two roughness configurations is dominated by the second eigen parameter. Although the first eigen parameter is the most influential for the variance of the terrain random field $h(x)$, it is not the most influential in terms of the variance of the resulting flow field. This is due to the particular shape of the second eigenvector and its impact on flow at the measurement location. Furthermore, except for input dimensions 1 and 2, no other input parameter has a considerable first-order effect on the distance. Using the approach described in Appendix B, we were further able to make an initial assessment of the interaction sensitivities and determine that significant interactions between inputs 1, 2, and 3 play an important role in the resulting turbulence intensity profiles. Using these insights, we were able to reduce the dimension of the KL random field to $d = 3$ and conduct a more rigorous study, further details of which are provided in [1].

7. Conclusions

This study investigates the application of active learning for global sensitivity analysis. We specifically propose a novel new learning function, termed the MUSIC (minimize uncertainty in Sobol index convergence) function, that leverages the main effect variances in Gaussian

Process-based Sensitivity Indices, i.e. the numerator in Sobol index definition. We compare this strategy with a more traditional active learning strategy aimed at accurately capturing output variance σ_Y^2 , i.e. the denominator of the Sobol index, using the existing EIGF and VIGF learning functions. Both strategies are further compared with simple sequential random sampling. We provide a careful analysis of the convergence for four analytical test functions and study the convergence behavior. An important insight that arises is that convergence in Sobol indices are difficult for active learning because they are defined as a ratio and the learning strategies must focus on either the numerator or the denominator and achieving rapid convergence in either, or both, is not sufficient to ensure convergence in the ratio.

Summarizing, we can make the following observations:

- In general, the VIGF function outperforms the EIGF, as does the MUSIC with VIGF when compared to the MUSIC with EIGF.
- The MUSIC learning functions are very successful at rapidly identifying small Sobol indices.
- For low-dimensional problems, the MUSIC with VIGF and distance pre-factor based on a Euclidean distance generally outperforms the other methods in terms of Sobol index convergence.
- The Euclidean distance-based pre-factor (D_2) is effective for low-dimensional problems but does not work well in high-dimensions.
- At best, the MUSIC learning strategy is comparable in performance to random sampling for high-dimensional problems.

Finally, we provide a motivating application where active learning is applied to learn global sensitivity indices for a set of large-scale boundary layer wind tunnel experiments. This demonstrates a real, practical use case and illustrates how it helped to significantly reduce experimental effort.

CRediT authorship contribution statement

Mohit S. Chauhan: Writing – original draft, Visualization, Validation, Software, Methodology, Investigation, Formal analysis. **Mariel Ojeda-Tuz:** Writing – review & editing, Investigation. **Ryan A. Catarelli:** Writing – review & editing, Supervision, Investigation. **Kurtis R. Gurley:** Writing – review & editing, Supervision, Project administration, Funding acquisition, Conceptualization. **Dimitrios Tsapetis:** Writing – review & editing, Software, Formal analysis. **Michael D. Shields:** Writing – review & editing, Supervision, Project administration, Methodology, Investigation, Funding acquisition, Conceptualization.

Declaration of competing interest

The authors declare the following financial interests/personal relationships which may be considered as potential competing interests: Michael Shields reports financial support was provided by National Science Foundation. Michael Shields reports financial support was provided by Defense Threat Reduction Agency.

Data availability

Data will be made available on request.

Acknowledgments

This material is based upon work supported by the National Science Foundation, United States under Grant Nos. 1930389 and 1930625. Dr. Dimitrios Tsapetis has been supported by Defense Threat Reduction Agency, United States, Award HDTRA1202000. Although results do not appear in the final paper, we are grateful to Prof. Pengfei Wei, who provided code for exploratory comparisons with the PVC method.

Appendix A

In this appendix, we show how to compute the mean and covariance functions for the main effect GPs under specific conditions. We then demonstrate that under conditions of widely used correlation models, these functions can be evaluated analytically without the need for numerical integrations.

A.1. Mean function of the main effect GPs

We begin by restating Eq. (22) as:

$$E[A(X_i)] = \int_{\mathbf{x} \sim i} \hat{y}(\mathbf{x}) \prod_{j \neq i} p_{X_j}(x_j) dx_j$$

Next, we simply plug in the expression for $\hat{y}$ from Eq. (14), which yields

$$E[A(X_i)] = \int_{\mathbf{x} \sim i} (\mathbf{f}(\mathbf{x})^T \boldsymbol{\beta} + \mathbf{r}(\mathbf{x})^T \mathbf{R}^{-1}(\mathbf{Y} - \mathbf{F}\boldsymbol{\beta})) \prod_{j \neq i} p_{X_j}(x_j) dx_j \quad (\text{A.1})$$

$$= \left[\int_{\mathbf{x} \sim i} \mathbf{f}(\mathbf{x})^T \prod_{j \neq i} p_{X_j}(x_j) dx_j \right] \boldsymbol{\beta} + \left[\int_{\mathbf{x} \sim i} \mathbf{r}(\mathbf{x})^T \prod_{j \neq i} p_{X_j}(x_j) dx_j \right] \times \mathbf{R}^{-1}(\mathbf{Y} - \mathbf{F}\boldsymbol{\beta}) \quad (\text{A.2})$$

Integrating the first term of Eq. (A.2) for a linear basis $\mathbf{f}(\mathbf{x}) = [1 \ x_1 \ x_2 \ \dots \ x_d]^T$.

$$\begin{aligned} \int_{\mathbf{x} \sim i} \mathbf{f}(\mathbf{x})^T \prod_{j \neq i} p_{X_j}(x_j) dx_j &= \int_{\mathbf{x} \sim i} [1 \ x_1 \ x_2 \ \dots \ x_d] \prod_{j \neq i} p_{X_j}(x_j) dx_j \\ &= [1 \ E[X_1] \ E[X_2] \ \dots \ x_i \ \dots \ E[X_d]] \end{aligned}$$

Integrating the second term of Eq. (A.2) yields:

$$\begin{aligned} \int_{\mathbf{x} \sim i} \mathbf{r}(\mathbf{x})^T \prod_{j \neq i} p_{X_j}(x_j) dx_j &= \int_{\mathbf{x} \sim i} \begin{bmatrix} \mathbf{R}(\mathbf{x}, \mathbf{x}^{(1)}) \\ \mathbf{R}(\mathbf{x}, \mathbf{x}^{(2)}) \\ \vdots \\ \mathbf{R}(\mathbf{x}, \mathbf{x}^{(n)}) \end{bmatrix} \prod_{j \neq i} p_{X_j}(x_j) dx_j \\ &= \int_{\mathbf{x} \sim i} \begin{bmatrix} \prod_k r_k(x_k, x_k^{(1)}) \\ \prod_k r_k(x_k, x_k^{(2)}) \\ \vdots \\ \prod_k r_k(x_k, x_k^{(n)}) \end{bmatrix} \prod_{j \neq i} p_{X_j}(x_j) dx_j \\ &= \begin{bmatrix} r_i(x_i, x_i^{(1)}) \prod_{j \neq i} \int_{x_j} r_j(x_j, x_j^{(1)}) p_{X_j}(x_j) dx_j \\ r_i(x_i, x_i^{(2)}) \prod_{j \neq i} \int_{x_j} r_j(x_j, x_j^{(2)}) p_{X_j}(x_j) dx_j \\ \vdots \\ r_i(x_i, x_i^{(n)}) \prod_{j \neq i} \int_{x_j} r_j(x_j, x_j^{(n)}) p_{X_j}(x_j) dx_j \end{bmatrix} \end{aligned}$$

where we recall that $\mathbf{x}^{(l)}, l = 1, \dots, n$ are the training samples and $r_k(x_k, x_k^{(l)})$ is the k th univariate correlation function expressed in Eq. (11). Finally, Eq. (A.2) can be expressed as:

$$\begin{aligned} E[A(X_i)] &= \begin{bmatrix} 1 \\ E[X_1] \\ E[X_2] \\ \vdots \\ x_i \\ \vdots \\ E[X_d] \end{bmatrix}^T \boldsymbol{\beta} + \begin{bmatrix} r_i(x_i, x_i^{(1)}) \prod_{j \neq i} \int_{x_j} r_j(x_j, x_j^{(1)}) p_{X_j}(x_j) dx_j \\ r_i(x_i, x_i^{(2)}) \prod_{j \neq i} \int_{x_j} r_j(x_j, x_j^{(2)}) p_{X_j}(x_j) dx_j \\ \vdots \\ r_i(x_i, x_i^{(n)}) \prod_{j \neq i} \int_{x_j} r_j(x_j, x_j^{(n)}) p_{X_j}(x_j) dx_j \end{bmatrix}^T \\ &\quad \times \mathbf{R}^{-1}(\mathbf{Y} - \mathbf{F}\boldsymbol{\beta}) \quad (\text{A.3}) \end{aligned}$$

Finally, in Appendix A.3, we show that the product of 1D integrals can be computed analytically for the Gaussian correlation model, which results in the following closed form for the mean of the main effect GP for uniformly distributed X :

$$\begin{aligned} E[A(X_i)] &= [1 \ E[X_1] \ E[X_2] \ \dots \ x_i \ \dots \ E[X_d]] \boldsymbol{\beta} \\ &\quad + \begin{bmatrix} r_i(x_i, x_i^{(1)}) \prod_{j \neq i} \frac{1}{b_j - a_j} \sqrt{\frac{\pi}{\theta_j}} [\Phi(\sqrt{2\theta_j}(b_j - x_j^{(1)})) - \Phi(\sqrt{2\theta_j}(a_j - x_j^{(1)}))] \\ r_i(x_i, x_i^{(2)}) \prod_{j \neq i} \frac{1}{b_j - a_j} \sqrt{\frac{\pi}{\theta_j}} [\Phi(\sqrt{2\theta_j}(b_j - x_j^{(2)})) - \Phi(\sqrt{2\theta_j}(a_j - x_j^{(2)}))] \\ \vdots \\ r_i(x_i, x_i^{(n)}) \prod_{j \neq i} \frac{1}{b_j - a_j} \sqrt{\frac{\pi}{\theta_j}} [\Phi(\sqrt{2\theta_j}(b_j - x_j^{(n)})) - \Phi(\sqrt{2\theta_j}(a_j - x_j^{(n)}))] \end{bmatrix}^T \\ &\quad \times \mathbf{R}^{-1}(\mathbf{Y} - \mathbf{F}\boldsymbol{\beta}) \quad (\text{A.4}) \end{aligned}$$

Note that in cases where the GP is trained with noisy data the above equation can be easily modified to incorporate the noise. We simply substitute the correlation matrix, $\mathbf{R}$, in the above equation by $\mathbf{R}_n$ defined as:

$$\mathbf{R}_n = ((1 - \tau)\mathbf{R} + \tau\mathbf{I})$$

where $\tau = \sigma_z^2 / \sigma_z^2 + \sigma_e^2$ and σ_e^2 is the variance due to the noise in the data and σ_z^2 is the total variance of the GP.

A.2. Covariance function of the main effect GPs

We again begin by restating the expression for the covariance function of the main effect GPs from Eq. (23):

$$\begin{aligned} \text{Cov}(A(X_{1i}), A(X_{2i})) &= \int_{\mathbf{x}_{1 \sim i}} \int_{\mathbf{x}_{2 \sim i}} \text{Cov}(\mathcal{Y}(\mathbf{x}_1), \mathcal{Y}(\mathbf{x}_2)) \prod_{j \neq i} p_{X_j}(x_{1j}) dx_{1j} \\ &\quad \times \prod_{k \neq i} p_{X_k}(x_{2k}) dx_{2k} \end{aligned}$$

The covariance of the GP $\mathcal{Y}(\mathbf{x})$ can be expressed as:

$$\begin{aligned} \text{Cov}(\mathcal{Y}(\mathbf{x}_1), \mathcal{Y}(\mathbf{x}_2)) &= \sigma_z^2 (\mathbf{r}(\mathbf{x}_1, \mathbf{x}_2) - \mathbf{r}(\mathbf{x}_1)^T \mathbf{R}^{-1} \mathbf{r}(\mathbf{x}_2) + \mathbf{t}(\mathbf{x}_1)^T \\ &\quad \times (\mathbf{F}^T \mathbf{R}^{-1} \mathbf{F})^{-1} \mathbf{t}(\mathbf{x}_2)), \end{aligned}$$

which yields the following covariance for the main effect GP:

$$\begin{aligned} \text{Cov}(A(X_{1i}), A(X_{2i})) &= \int_{\mathbf{x}_{1 \sim i}} \int_{\mathbf{x}_{2 \sim i}} \left[\sigma_z^2 (\mathbf{r}(\mathbf{x}_1, \mathbf{x}_2) - \mathbf{r}(\mathbf{x}_1)^T \mathbf{R}^{-1} \mathbf{r}(\mathbf{x}_2) + \mathbf{t}(\mathbf{x}_1)^T (\mathbf{F}^T \mathbf{R}^{-1} \mathbf{F})^{-1} \mathbf{t}(\mathbf{x}_2)) \right] \\ &\quad \times \prod_{j \neq i} p_{X_j}(x_{1j}) \prod_{k \neq i} p_{X_k}(x_{2k}) dx_{2k} \quad (\text{A.5}) \end{aligned}$$

Expanding this, we can express the first term of Eq. (A.5) as:

$$\begin{aligned} \sigma_z^2 \int_{\mathbf{x}_{1 \sim i}} \int_{\mathbf{x}_{2 \sim i}} \mathbf{r}(\mathbf{x}_1, \mathbf{x}_2) \prod_{j \neq i} p_{X_j}(x_{1j}) \prod_{k \neq i} p_{X_k}(x_{2k}) dx_{2k} \\ &= \sigma_z^2 \int_{\mathbf{x}_{1 \sim i}} \int_{\mathbf{x}_{2 \sim i}} \prod_{l=1}^d r_l(x_{1l}, x_{2l}) \prod_{j \neq i} p_{X_j}(x_{1j}) \prod_{k \neq i} p_{X_k}(x_{2k}) dx_{2k} \\ &= \sigma_z^2 r_i(x_{1i}, x_{2i}) \prod_{l \neq i} \int_{\mathbf{x}_{1 \sim i}} \int_{\mathbf{x}_{2 \sim i}} r_l(x_{1l}, x_{2l}) p_{X_l}(x_{1l}) p_{X_l}(x_{2l}) dx_{2l} \end{aligned}$$

$$\int_{\mathbf{x}_{1:i}} \mathbf{t}(\mathbf{x})^T \prod_{j \neq i} p_{X_j}(x_j) dx_j = \mathbf{F}^T \mathbf{R}^{-1} \begin{bmatrix} r_i(x_i, x_i^{(1)}) \prod_{j \neq i} \frac{1}{b_j - a_j} \sqrt{\frac{\pi}{\theta_j}} [\Phi(\sqrt{2\theta_j}(b_j - x_j^{(1)})) - \Phi(\sqrt{2\theta_j}(a_j - x_j^{(1)}))] \\ r_i(x_i, x_i^{(2)}) \prod_{j \neq i} \frac{1}{b_j - a_j} \sqrt{\frac{\pi}{\theta_j}} [\Phi(\sqrt{2\theta_j}(b_j - x_j^{(2)})) - \Phi(\sqrt{2\theta_j}(a_j - x_j^{(2)}))] \\ \vdots \\ r_i(x_i, x_i^{(n)}) \prod_{j \neq i} \frac{1}{b_j - a_j} \sqrt{\frac{\pi}{\theta_j}} [\Phi(\sqrt{2\theta_j}(b_j - x_j^{(n)})) - \Phi(\sqrt{2\theta_j}(a_j - x_j^{(n)}))] \end{bmatrix} - \begin{bmatrix} 1 \\ E[X_1] \\ E[X_2] \\ \vdots \\ x_i \\ \vdots \\ E[X_d] \end{bmatrix}$$

Box 1.

In [Appendix A.3](#), we derive the closed form for this integral given a uniform distribution and Gaussian correlation model, which yields:

$$\begin{aligned} \sigma_z^2 \int_{\mathbf{x}_{1:i}} \int_{\mathbf{x}_{2:i}} r(\mathbf{x}_1, \mathbf{x}_2) \prod_{j \neq i} p_{X_j}(x_{1j}) dx_{1j} \prod_{k \neq i} p_{X_k}(x_{2k}) dx_{2k} \\ = \sigma_z^2 r_i(x_{1i}, x_{2i}) \prod_{l \neq i} \left(\frac{1}{b_l - a_l} \right)^2 \sqrt{\frac{\pi}{\theta_l}} \\ \times \left[a_l + b_l - 2[a_l \Phi(\sqrt{2\theta_l}(b_l - a_l)) + b_l \Phi(\sqrt{2\theta_l}(a_l - b_l))] \right. \\ \left. - \frac{1}{\sqrt{\pi\theta_l}} [1 - \exp\{-\theta_l(b_l - a_l)^2\}] \right] \end{aligned}$$

Next, the second term of the Eq. (A.5) is:

$$\begin{aligned} -\sigma_z^2 \int_{\mathbf{x}_{1:i}} \int_{\mathbf{x}_{2:i}} \mathbf{r}(\mathbf{x}_1)^T \mathbf{R}^{-1} \mathbf{r}(\mathbf{x}_2) \prod_{j \neq i} p_{X_j}(x_{1j}) dx_{1j} \prod_{k \neq i} p_{X_k}(x_{2k}) dx_{2k} \\ = -\sigma_z^2 \left[\int_{\mathbf{x}_{1:i}} \mathbf{r}(\mathbf{x}_1)^T \prod_{j \neq i} p_{X_j}(x_{1j}) dx_{1j} \right] \mathbf{R}^{-1} \\ \times \left[\int_{\mathbf{x}_{2:i}} \mathbf{r}(\mathbf{x}_2) \prod_{k \neq i} p_{X_k}(x_{2k}) dx_{2k} \right] \end{aligned}$$

where the integrals in brackets can be solved in closed form as derived in [Appendix A.1](#). Finally, the third term of Eq. (A.5) is:

$$\begin{aligned} -\sigma_z^2 \int_{\mathbf{x}_{1:i}} \int_{\mathbf{x}_{2:i}} \mathbf{t}(\mathbf{x}_1)^T (\mathbf{F}^T \mathbf{R}^{-1} \mathbf{F})^{-1} \mathbf{t}(\mathbf{x}_2) \prod_{j \neq i} p_{X_j}(x_{1j}) dx_{1j} \prod_{j \neq i} p_{X_j}(x_{2j}) dx_{2j} \\ = -\sigma_z^2 \left[\int_{\mathbf{x}_{1:i}} \mathbf{t}(\mathbf{x}_1)^T \prod_{j \neq i} p_{X_j}(x_{1j}) dx_{1j} \right] (\mathbf{F}^T \mathbf{R}^{-1} \mathbf{F})^{-1} \\ \times \left[\int_{\mathbf{x}_{2:i}} \mathbf{t}(\mathbf{x}_2) \prod_{j \neq i} p_{X_j}(x_{2j}) dx_{2j} \right] \end{aligned}$$

where, $\hat{\sigma}_z^2$ is the variance of the GP, $\mathcal{Y}$, and $\mathbf{t}(\mathbf{x}) = \mathbf{F}^T \mathbf{R}^{-1} \mathbf{r}(\mathbf{x}) - \mathbf{f}(\mathbf{x})$ as defined in Eq. (16). Finally, the integral of $\mathbf{t}(\mathbf{x})$ over the joint distribution in brackets can be easily obtained from the equations in [Appendix A.1](#), as follows

$$\begin{aligned} \int_{\mathbf{x}_{1:i}} \mathbf{t}(\mathbf{x})^T \prod_{j \neq i} p_{X_j}(x_j) dx_j \\ = \int_{\mathbf{x}_{1:i}} \left[\mathbf{F}^T \mathbf{R}^{-1} \mathbf{r}(\mathbf{x}) - \mathbf{f}(\mathbf{x}) \right] \prod_{j \neq i} p_{X_j}(x_j) dx_j \\ = \mathbf{F}^T \mathbf{R}^{-1} \int_{\mathbf{x}_{1:i}} \mathbf{r}(\mathbf{x}) \prod_{j \neq i} p_{X_j}(x_j) dx_j - \int_{\mathbf{x}_{1:i}} \mathbf{f}(\mathbf{x}) \prod_{j \neq i} p_{X_j}(x_j) dx_j \\ = \mathbf{F}^T \mathbf{R}^{-1} \begin{bmatrix} r_i(x_i, x_i^{(1)}) \prod_{j \neq i} \int_{x_j} r_j(x_j, x_j^{(1)}) p_{X_j}(x_j) dx_j \\ r_i(x_i, x_i^{(2)}) \prod_{j \neq i} \int_{x_j} r_j(x_j, x_j^{(2)}) p_{X_j}(x_j) dx_j \\ \vdots \\ r_i(x_i, x_i^{(n)}) \prod_{j \neq i} \int_{x_j} r_j(x_j, x_j^{(n)}) p_{X_j}(x_j) dx_j \end{bmatrix} - \begin{bmatrix} 1 \\ E[X_1] \\ E[X_2] \\ \vdots \\ x_i \\ \vdots \\ E[X_d] \end{bmatrix} \end{aligned}$$

The above expression can be further simplified for a given uniform distribution and Gaussian correlation model as: $\int_{\mathbf{x}_{1:i}} \mathbf{t}(\mathbf{x})^T \prod_{j \neq i} p_{X_j}(x_j) dx_j$ is given in [Box 1](#)

Thus, the main effect GP, $A(X_i)$, can be completely defined by the mean function, $E[A(X_i)]$, and covariance function, $\text{Cov}(A(X_{1i}), A(X_{2i}))$, derived in [Appendices A.1](#) and [A.2](#).

Again in case of noisy training data, we simply substitute the correlation matrix, $\mathbf{R}$, in the above equation by $\mathbf{R}_n$ as defined below:

$$\mathbf{R}_n = ((1 - \tau)\mathbf{R} + \tau\mathbf{I})$$

where, $\tau = \sigma_e^2 / \sigma_z^2$ and σ_e^2 is the variance due to the noise in the data and σ_z^2 is the total variance of the GP.

A.3. Closed-form integration of the Gaussian correlation model

Let x_1 and x_2 be univariate sample points drawn from a uniform distribution over the range (a, b) , $X_1, X_2 \sim U(a, b)$, having Gaussian correlation (from Eq. (11)) given by:

$$r(x_1, x_2 | \theta) = \exp\{-\theta(x_1 - x_2)^2\}$$

Note that, without loss of generality, the arbitrary sample points can be transformed to a uniform distribution through a suitable isoprobabilistic transformation (e.g. Nataf [53]). As discussed in the previous appendices, we need to integrate this correlation model over the pdf $p_{X_1}(x_1)$ as follows:

$$\begin{aligned} \int_{-\infty}^{\infty} r(\theta, x_1, x_2) p_{X_1}(x_1) dx_1 &= \int_a^b \exp\{-\theta(x_1 - x_2)^2\} \frac{1}{b-a} dx_1 \\ &= \frac{1}{b-a} \int_a^b \exp\{-\theta(x_1 - x_2)^2\} dx_1 \\ &= \frac{1}{b-a} \sqrt{2\pi(1/2\theta)} \int_a^b \frac{1}{\sqrt{2\pi(1/2\theta)}} \\ &\quad \times \exp\left\{-\frac{1}{2} \frac{(x_1 - x_2)^2}{(1/2\theta)}\right\} dx_1 \\ &= \frac{1}{b-a} \sqrt{\frac{\pi}{\theta}} \left[\Phi\left(\sqrt{2\theta}(b - x_2)\right) \right. \\ &\quad \left. - \Phi\left(\sqrt{2\theta}(a - x_2)\right) \right] \end{aligned}$$

where, $\Phi(\cdot)$ is the standard normal cumulative distribution function. Further, we need to integrate the above result with respect to x_2 as

$$\begin{aligned} \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} r(\theta, x_1, x_2) p_{X_1}(x_1) p_{X_2}(x_2) dx_1 dx_2 \\ = \int_{-\infty}^{\infty} \left[\frac{1}{b-a} \sqrt{\frac{\pi}{\theta}} \left[\Phi\left(\sqrt{2\theta}(b - x_2)\right) \right. \right. \\ \left. \left. - \Phi\left(\sqrt{2\theta}(a - x_2)\right) \right] \right] p_{X_2}(x_2) dx_2 \\ = \left(\frac{1}{b-a} \right)^2 \sqrt{\frac{\pi}{\theta}} \int_a^b \left[\Phi\left(\sqrt{2\theta}(b - x_2)\right) \right. \\ \left. - \Phi\left(\sqrt{2\theta}(a - x_2)\right) \right] dx_2 \end{aligned} \quad (\text{A.6})$$

We see that the above integral has two identical components that differ only by the constant a or b . Let us introduce a dummy variable, z , which

can take values a or b . Integration by parts yields

$$\begin{aligned} & \int_a^b (1) \times \Phi(\sqrt{2\theta}(z - x_2)) dx_2 \\ &= \Phi(\sqrt{2\theta}(z - s)) x_2 \Big|_a^b - \int_a^b x_2 \phi(\sqrt{2\theta}(z - x_2)) (-\sqrt{2\theta}) dx_2 \end{aligned}$$

where, $\phi(\cdot)$ is the standard normal probability density function. In this integration by parts and with the use of a change of variables as $t = 2\theta(z - x_2)^2$, the second term can be expressed as:

$$\begin{aligned} & \int_a^b x_2 \phi(\sqrt{2\theta}(z - x_2)) (-\sqrt{2\theta}) dx_2 \\ &= \int_a^b [z - (z - x_2)] (-\sqrt{2\theta}) \frac{1}{\sqrt{2\pi}} \exp\left\{-\frac{1}{2} \frac{(z - x_2)^2}{(1/2\theta)}\right\} dx_2 \\ &= -z \int_a^b \frac{1}{\sqrt{2\pi(1/2\theta)}} \exp\left\{-\frac{1}{2} \frac{(z - x_2)^2}{(1/2\theta)}\right\} dx_2 \\ &\quad + \int_a^b \frac{(z - x_2)}{(1/\sqrt{2\theta})} \frac{1}{\sqrt{2\pi}} \exp\left\{-\frac{1}{2} \frac{(z - x_2)^2}{(1/2\theta)}\right\} dx_2 \\ &= -z \Phi(\sqrt{2\theta}(x_2 - z)) \Big|_a^b + \int_{2\theta(z-a)^2}^{2\theta(z-b)^2} \sqrt{t} \frac{1}{\sqrt{2\pi}} \exp\left\{-\frac{1}{2} t\right\} \\ &\quad \times \frac{dt}{2\sqrt{t}(-\sqrt{2\theta})} \\ &= -z \Phi(\sqrt{2\theta}(x_2 - z)) \Big|_a^b - \frac{1}{4\sqrt{\pi\theta}} \int_{2\theta(z-a)^2}^{2\theta(z-b)^2} \exp\left\{-\frac{1}{2} t\right\} dt \\ &= -z \Phi(\sqrt{2\theta}(x_2 - z)) \Big|_a^b + \frac{1}{2\sqrt{\pi\theta}} \exp\left\{-\frac{1}{2} t\right\} \Big|_{2\theta(z-a)^2}^{2\theta(z-b)^2} \end{aligned}$$

Using this expression the integral in Eq. (A.6), can be evaluate as:

$$\begin{aligned} & \int_a^b \Phi(\sqrt{2\theta}(b - x_2)) dx_2 - \int_a^b \Phi(\sqrt{2\theta}(a - x_2)) dx_2 \\ &= a + b - 2[a\Phi(\sqrt{2\theta}(b - a)) + b\Phi(\sqrt{2\theta}(a - b))] \\ &\quad - \frac{1}{\sqrt{\pi\theta}} [1 - \exp\{-\theta(b - a)^2\}] \end{aligned}$$

Finally, we can compute the double integration of the correlation function in closed form as:

$$\begin{aligned} & \int_{-\infty}^{\infty} \int_{-\infty}^{\infty} r(\theta, x_1, x_2) p_{X_1}(x_1) p_{X_2}(x_2) dx_1 dx_2 \\ &= \left(\frac{1}{b-a}\right)^2 \sqrt{\frac{\pi}{\theta}} \left[a + b - 2[a\Phi(\sqrt{2\theta}(b - a)) + b\Phi(\sqrt{2\theta}(a - b))] \right. \\ &\quad \left. - \frac{1}{\sqrt{\pi\theta}} [1 - \exp\{-\theta(b - a)^2\}] \right] \end{aligned} \quad (A.7)$$

Appendix B. Interaction sensitivities

The interaction sensitivities are computed using a similar procedure as the main effect sensitivities. These indices can be expressed using following equation

$$\tilde{S}_p(\omega) = \frac{\text{Var}_{X_p}(E_{X_{\sim p}}[\mathcal{Y}(X, \omega) | X_p])}{E[\text{Var}(\mathcal{Y}(X, \omega))]} \quad (B.1)$$

where, p is a subset of indices of the input dimensions(i.e. $p \in \{1, 2, \dots, d\}$).

To compute the interaction Sobol indices, we start by taking the expectation of the GP $\mathcal{Y}(X, \omega)$ over all the inputs except ' X_p ' and denote this as

$$A(X_p, \omega) = E_{X_{\sim p}}[\mathcal{Y}(X, \omega) | X_p] \quad (B.2)$$

Since, $\mathcal{Y}(X, \omega)$ is a Gaussian process and expectation is a linear operator, $A(X_p, \omega)$ is also a GP referred to as the interaction effect GP.

The mean and covariance function of the interaction effect GP can be determined by integrating the original GP with respect to the joint

probability measure over all inputs except X_p . Considering independent inputs, the mean function is given by [45]:

$$\mu_{A_p}(x_p) = \mathbb{E}[A(X_p)] = \int_{x_{\sim p}} \hat{y}(x) \prod_{j \notin p} p_{X_j}(x_j) dx_j \quad (B.3)$$

which is easily computed as a product of one-dimensional integrals. Again considering independent inputs, the covariance function of the interaction effect GP is given by [45]:

$$\begin{aligned} & \text{Cov}(A(X_{1p}), A(X_{2p})) \\ &= \int_{x_{1\sim p}} \int_{x_{2\sim p}} \text{Cov}(\mathcal{Y}(x_1), \mathcal{Y}(x_2)) \prod_{j \notin p} p_{X_j}(x_{1j}) dx_{1j} \prod_{j \notin p} p_{X_j}(x_{2j}) dx_{2j} \end{aligned} \quad (B.4)$$

which can be expressed as a product of 2-dimensional integrals. After obtaining the mean and covariance function of the interaction effect GP, the estimates for sensitivities indices are obtained in a similar procedure as main effect sensitivities. Note that the contributions of the interaction sensitivities are not used during the learning process.

References

- [1] Shields Michael D, Gurley Kurtis, Catarelli Ryan, Chauhan Mohit, Ojeda-Tuz Mariel, Masters Forrest. Active learning applied to automated physical systems increases the rate of discovery. *Sci Rep* 2023.
- [2] Shahriari Bobak, Swersky Kevin, Wang Ziyu, Adams Ryan P, De Freitas Nando. Taking the human out of the loop: A review of Bayesian optimization. *Proc IEEE* 2015;104(1):148–75.
- [3] Frazier Peter I. Bayesian optimization. In: *Recent advances in optimization and modeling of contemporary problems*. Inform; 2018, p. 255–78.
- [4] Greenhill Stewart, Rana Santu, Gupta Sunil, Vellanki Pratibha, Venkatesh Svetha. Bayesian optimization for adaptive experimental design: A review. *IEEE Access* 2020;8:13937–48.
- [5] Moćkus Jonas. On Bayesian methods for seeking the extremum. In: *Optimization techniques IFIP technical conference: Novosibirsk, July 1–7, 1974*. Springer; 1975, p. 400–4.
- [6] Jones Donald R, Schonlau Matthias, Welch William J. Efficient global optimization of expensive black-box functions. *J Glob Optim* 1998;13(4):455.
- [7] Huang Deng, Allen Theodore T, Notz William I, Zeng Ning, et al. Global optimization of stochastic black-box systems via sequential Kriging meta-models. *J Glob Optim* 2006;34(3):441–66.
- [8] Frazier Peter, Powell Warren, Dayanik Savas. The knowledge-gradient policy for correlated normal beliefs. *INFORMS J Comput* 2009;21(4):599–613.
- [9] Scott Warren, Frazier Peter, Powell Warren. The correlated knowledge gradient for simulation optimization of continuous parameters using gaussian process regression. *SIAM J Optim* 2011;21(3):996–1026.
- [10] Hennig Philipp, Schuler Christian J. Entropy search for information-efficient global optimization. *J Mach Learn Res* 2012;13(6).
- [11] Hernández-Lobato José Miguel, Hoffman Matthew W, Ghahramani Zoubin. Predictive entropy search for efficient global optimization of black-box functions. *Adv Neural Inf Process Syst* 2014;27.
- [12] Bichon Barron J, Eldred Michael S, Swiler Laura Painton, Mahadevan Sandaran, McFarland John M. Efficient global reliability analysis for nonlinear implicit performance functions. *AIAA J* 2008;46(10):2459–68.
- [13] Echard Benjamin, Gayton Nicolas, Lemaire Maurice. AK-MCS: an active learning reliability method combining Kriging and Monte Carlo simulation. *Struct Saf* 2011;33(2):145–54.
- [14] Zhang Jinhao, Xiao Mi, Gao Liang. An active learning reliability method combining Kriging constructed with exploration and exploitation of failure region and subset simulation. *Reliab Eng Syst Saf* 2019;188:90–102.
- [15] Xu Chunlong, Chen Weidong, Ma Jingxin, Shi Yaqin, Lu Shengzhuo. AK-MSS: An adaptation of the AK-MCS method for small failure probabilities. *Struct Saf* 2020;86:101971.
- [16] Sundar VS, Shields Michael D. Reliability analysis using adaptive Kriging surrogates with multimodel inference. *ASCE-ASME J Risk Uncertain Eng Syst A* 2019;5(2):04019004.
- [17] Razaaly Nassim, Congedo Pietro Marco. Extension of AK-MCS for the efficient computation of very small failure probabilities. *Reliab Eng Syst Saf* 2020;203:107084.
- [18] El Haj Abdul-Kader, Soubra Abdul-Hamid. Improved active learning probabilistic approach for the computation of failure probability. *Struct Saf* 2021;88:102011.
- [19] Wei Pengfei, Zheng Yu, Fu Jiangfeng, Xu Yuannan, Gao Weikai. An expected integrated error reduction function for accelerating Bayesian active learning of failure probability. *Reliab Eng Syst Saf* 2023;231:108971.
- [20] Wang Jinsheng, Xu Guojie, Yuan Peng, Li Yongle, Kareem Ahsan. An efficient and versatile Kriging-based active learning method for structural reliability analysis. *Reliab Eng Syst Saf* 2024;241:109670.

- [21] Lam Chen Quin. Sequential adaptive designs in computer experiments for response surface model fit (Ph.D. thesis), The Ohio State University; 2008.
- [22] Cheng Kai, Lu Zhenzhou. Active learning polynomial chaos expansion for reliability analysis by maximizing expected indicator function prediction error. *Internat J Numer Methods Engrg* 2020;121(14):3159–77.
- [23] Novák Lukáš, Shields Michael D, Sadílek Václav, Vořechovský Miroslav. Active learning-based domain adaptive localized polynomial chaos expansion. 2023, arXiv preprint arXiv:2301.13635.
- [24] Zhang Jian, Gong Weijie, Yue Xinxi, Shi Maolin, Chen Lei. Efficient reliability analysis using prediction-oriented active sparse polynomial chaos expansion. *Reliab Eng Syst Saf* 2022;228:108749.
- [25] Haut Juan Mario, Paoletti Mercedes E, Plaza Javier, Li Jun, Plaza Antonio. Active learning with convolutional neural networks for hyperspectral image classification using a new Bayesian approach. *IEEE Trans Geosci Remote Sens* 2018;56(11):6440–61.
- [26] Schröder Christopher, Niekler Andreas. A survey of active learning for text classification using deep neural networks. 2020, arXiv preprint arXiv:2008.07267.
- [27] Zhu Rong, Chen Yuan, Peng Weiwen, Ye Zhi-Sheng. Bayesian deep-learning for RUL prediction: An active learning perspective. *Reliab Eng Syst Saf* 2022;228:108758.
- [28] DeLoach Richard. Applications of modern experiment design to wind tunnel testing at nasa langley research center. In: 36th AIAA aerospace sciences meeting and exhibit. 1998, p. 713.
- [29] Hill Raymond R, Leggio Derek A, Capehart Shay R, Roesener August G. Examining improved experimental designs for wind tunnel testing using Monte Carlo sampling methods. *Qual Reliab Eng Int* 2011;27(6):795–803.
- [30] VanDercreek Colin P, Amiri-Simkooei Alireza, Snellen Mirjam, Ragni Daniele. Experimental design and stochastic modeling of hydrodynamic wave propagation within cavities for wind tunnel acoustic measurements. *Int J Aeroacoust* 2019;18(8):752–79.
- [31] Kusne A Gilad, Yu Heshan, Wu Changming, Zhang Huairuo, Hattrick-Simpers Jason, DeCost Brian, Sarker Suchismita, Osos Corey, Toher Cormac, Curtarolo Stefano, et al. On-the-fly closed-loop materials discovery via Bayesian active learning. *Nat Commun* 2020;11(1):5966.
- [32] Jablonka Kevin Maik, Jothiappan Giriprasad Melpatti, Wang Shefang, Smit Berend, Yoo Brian. Bias free multiobjective active learning for materials design and discovery. *Nat Commun* 2021;12(1):2312.
- [33] Sverchkov Yuriy, Craven Mark. A review of active learning approaches to experimental design for uncovering biological networks. *PLoS Comput Biol* 2017;13(6):e1005466.
- [34] Reker Daniel. Practical considerations for active machine learning in drug discovery. *Drug Discov Today: Technol* 2019;32–33:73–9. <http://dx.doi.org/10.1016/j.ddtec.2020.06.001>, Artificial Intelligence.
- [35] Osugi Thomas, Kim Deng, Scott Stephen. Balancing exploration and exploitation: A new algorithm for active machine learning. In: Fifth IEEE international conference on data mining. ICDM'05, IEEE; 2005, p. 8–pp.
- [36] Shan Songqing, Wang G Gary. Survey of modeling and optimization strategies to solve high-dimensional design problems with computationally-expensive black-box functions. *Struct Multidiscip Optim* 2010;41:219–41.
- [37] Sobol Ilya M. Sensitivity analysis for non-linear mathematical models. *Math Model Comput Exp* 1993;1:407–14.
- [38] Saltelli Andrea, Sobol Ilya M. About the use of rank transformation in sensitivity analysis of model output. *Reliab Eng Syst Saf* 1995;50(3):225–39.
- [39] Archer GEB, Saltelli Andrea, Sobol Ilya Meyerovich. Sensitivity measures, ANOVA-like techniques and the use of bootstrap. *J Stat Comput Simul* 1997;58(2):99–120.
- [40] Rasmussen Carl Edward, Williams Christopher KI, et al. Gaussian processes for machine learning, vol. 1, Springer; 2006.
- [41] Matheron Georges. Principles of geostatistics. *Econ Geol* 1963;58(8):1246–66.
- [42] Cressie Noel. Statistics for spatial data. John Wiley & Sons; 2015.
- [43] Martin Jay D, Simpson Timothy W. Use of Kriging models to approximate deterministic computer models. *AIAA J* 2005;43(4):853–63.
- [44] Marrel Amandine, Iooss Bertrand, Laurent Béatrice, Roustant Olivier. Calculations of Sobol indices for the Gaussian process metamodel. *Reliab Eng Syst Saf* 2009;94(3):742–51. <http://dx.doi.org/10.1016/j.ress.2008.07.008>.
- [45] Oakley J E, O'Hagan A. Probabilistic sensitivity analysis of complex models: a Bayesian approach. *J R Stat Soc Ser B Stat Methodol* 2004;66(3):751–69. <http://dx.doi.org/10.1111/j.1467-9868.2004.05304.x>.
- [46] Mohammadi Hossein, Challenor Peter. Sequential adaptive design for emulating costly computer codes. 2022, arXiv preprint arXiv:2206.12113.
- [47] Blanchard Antoine, Sapsis Themistoklis. Output-weighted optimal sampling for Bayesian experimental design and uncertainty quantification. *SIAM/ASA J Uncertain Quant* 2021;9(2):564–92.
- [48] Wei Pengfei, Zhang Xing, Beer Michael. Adaptive experiment design for probabilistic integration. *Comput Methods Appl Mech Engrg* 2020;365:113035.
- [49] Song Jingwen, Wei Pengfei, Valdebenito Marcos A, Faes Matthias, Beer Michael. Data-driven and active learning of variance-based sensitivity indices with Bayesian probabilistic integration. *Mech Syst Signal Process* 2022;163:108106.
- [50] Beck Joakim, Guillas Serge. Sequential design with mutual information for computer experiments (MICE): Emulation of a tsunami model. *SIAM/ASA J Uncertain Quant* 2016;4(1):739–66.
- [51] Aggarwal Charu C, Hinneburg Alexander, Keim Daniel A. On the surprising behavior of distance metrics in high dimensional space. In: Database theory—ICDT 2001: 8th international conference London, UK, January 4–6, 2001 proceedings 8. Springer; 2001, p. 420–34.
- [52] Catarelli Ryan A, Fernández-Cabán Pedro L, Phillips Brian M, Bridge Jennifer A, Masters Forrest J, Gurley Kurtis R, Prevatt David O. Automation and new capabilities in the university of florida NHERI boundary layer wind tunnel. *Front Built Environ* 2020;6:558151.
- [53] Lebrun Regis, Dutfoy Anne. An innovating analysis of the Nataf transformation from the copula viewpoint. *Probab Eng Mech* 2009;24(3):312–20.